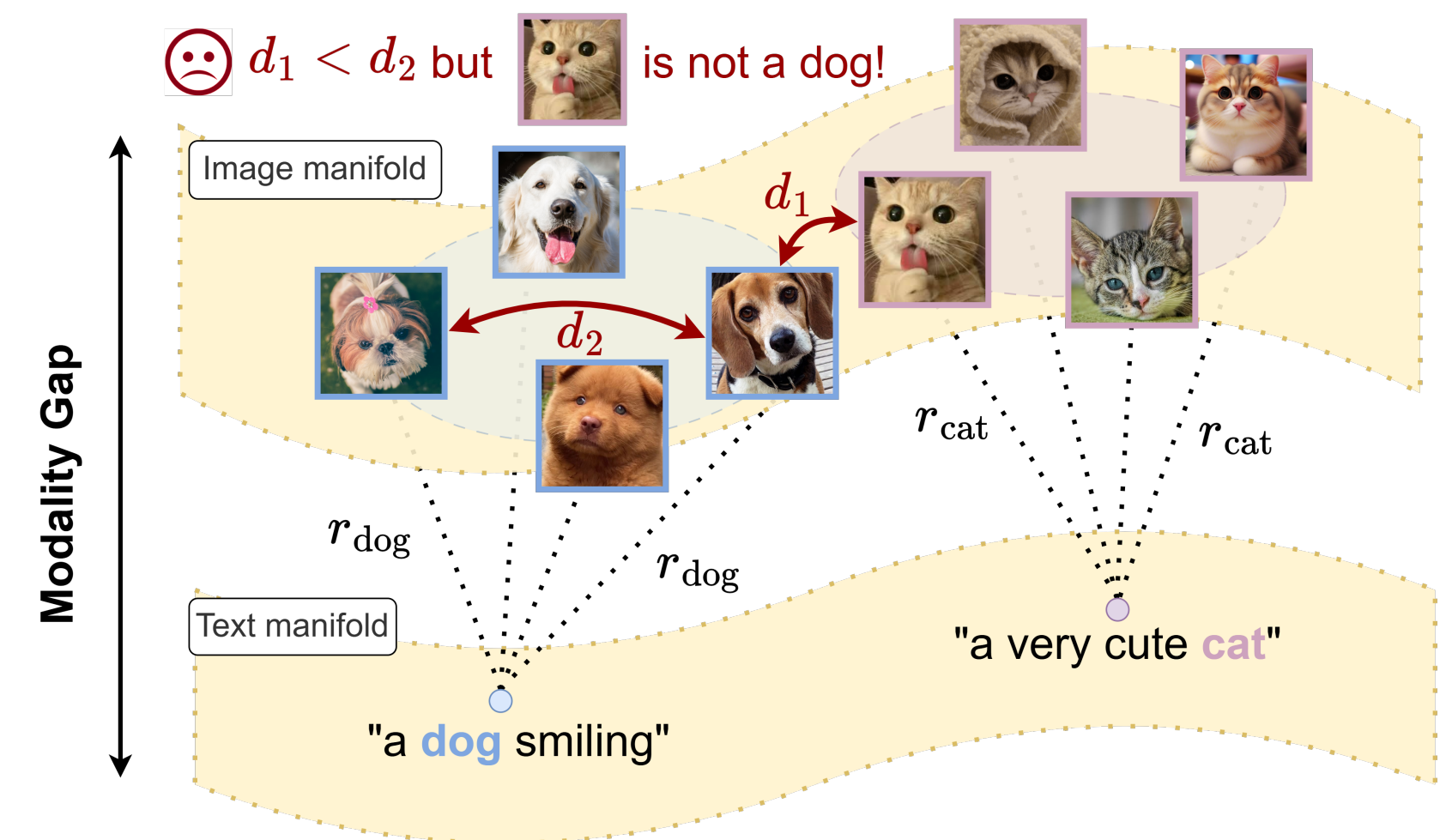
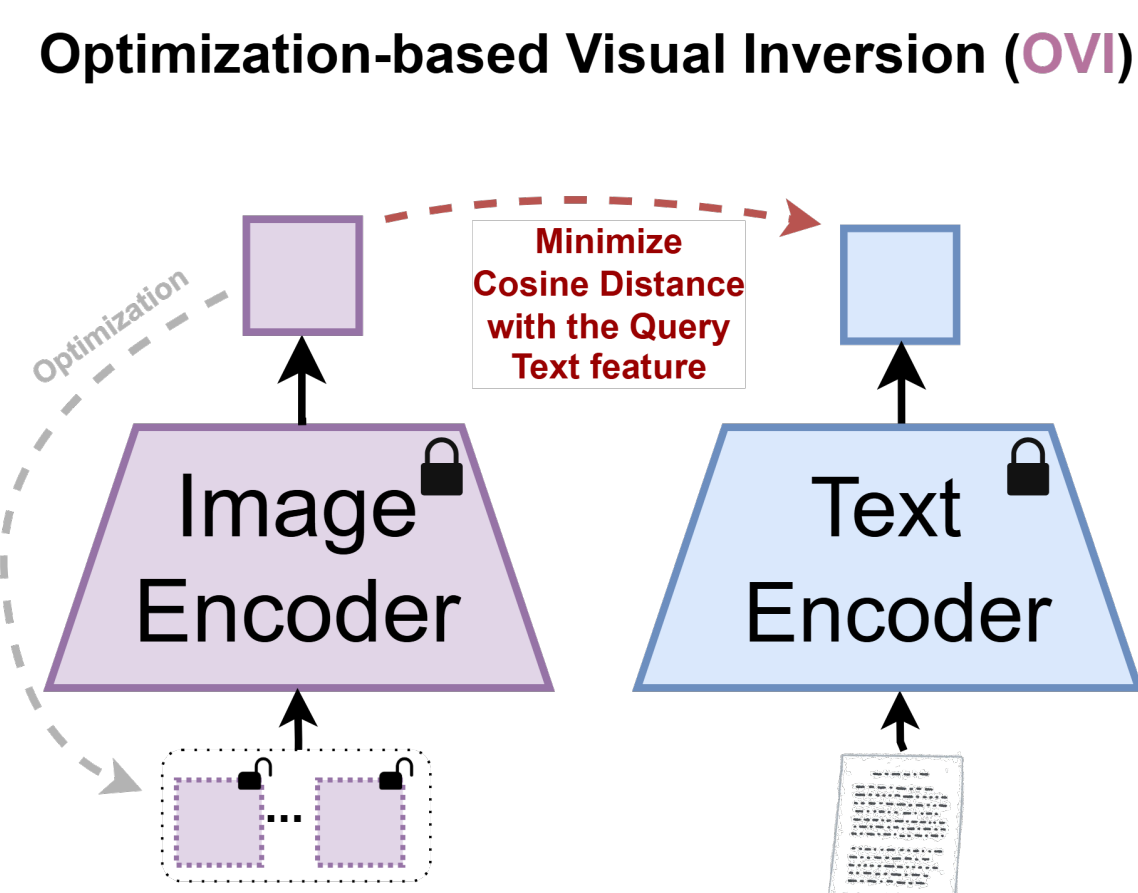
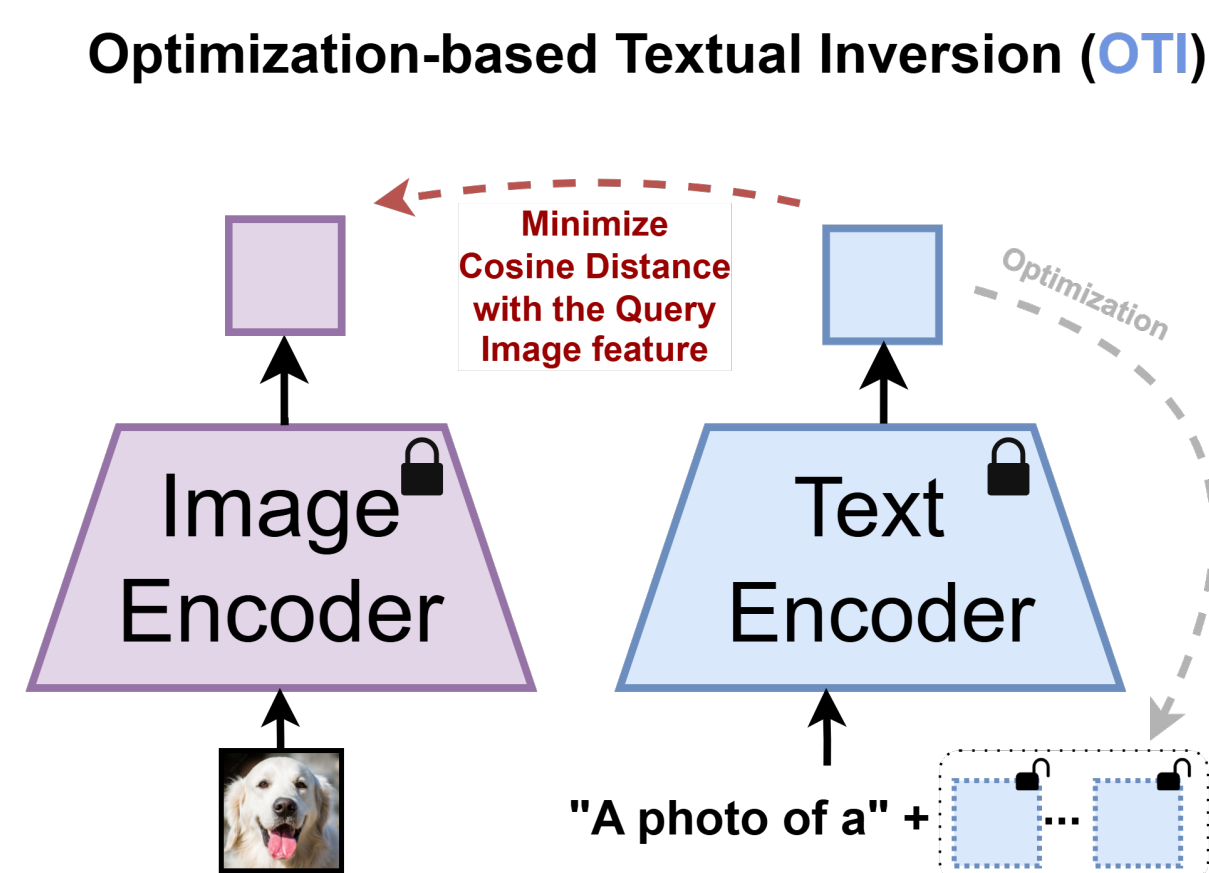




Cross the Gap: Exposing the Intra-modal Misalignment in CLIP via Modality Inversion

*Marco Mistretta, *Alberto Baldrati, *Lorenzo Agnolucci, Marco Bertini, Andrew D. Bagdanov



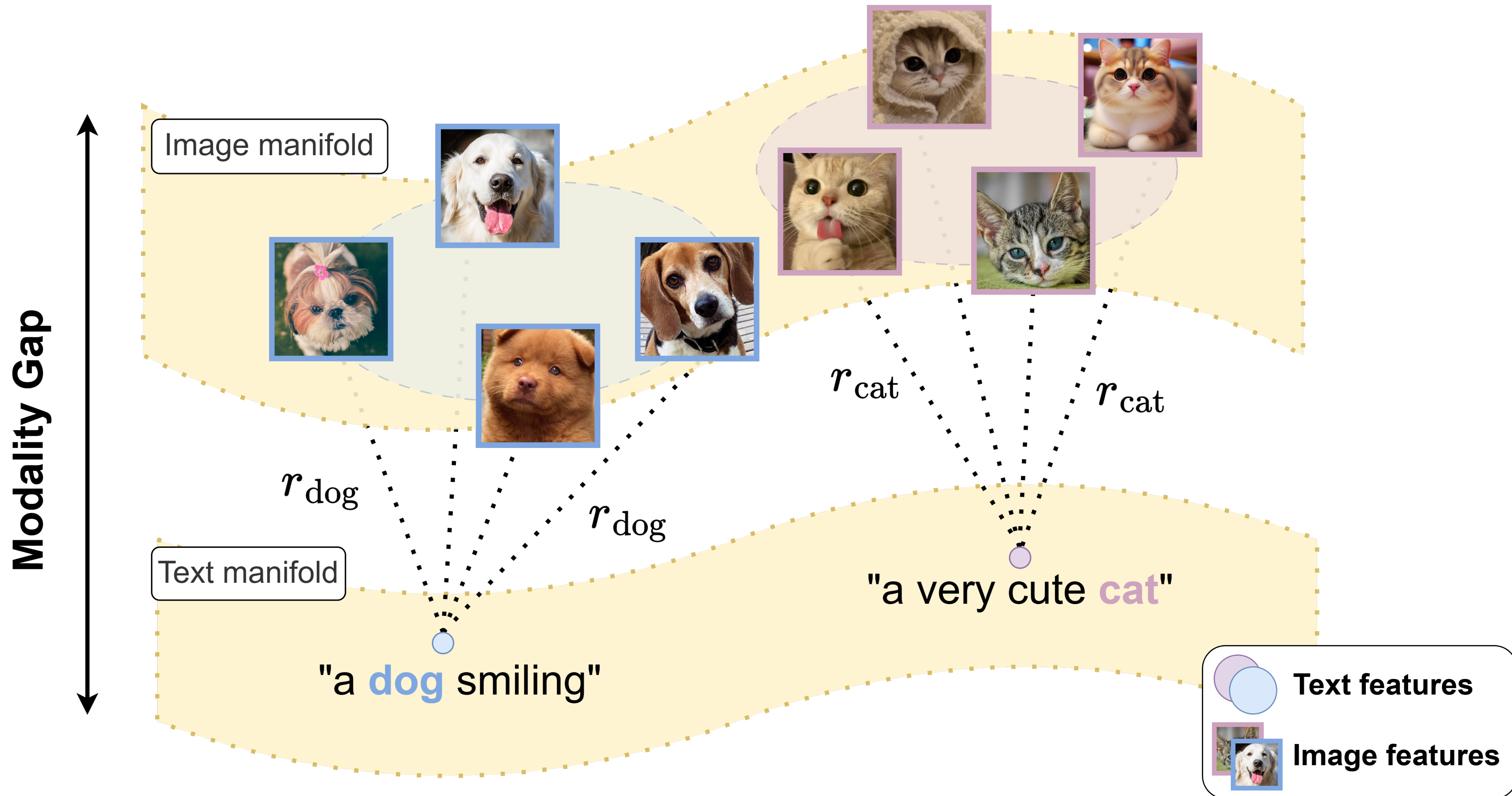
TLDR;

We adapt **OTI**, and we introduce **OVI**...

...in order to expose **CLIP Intra-modal Misalignment**

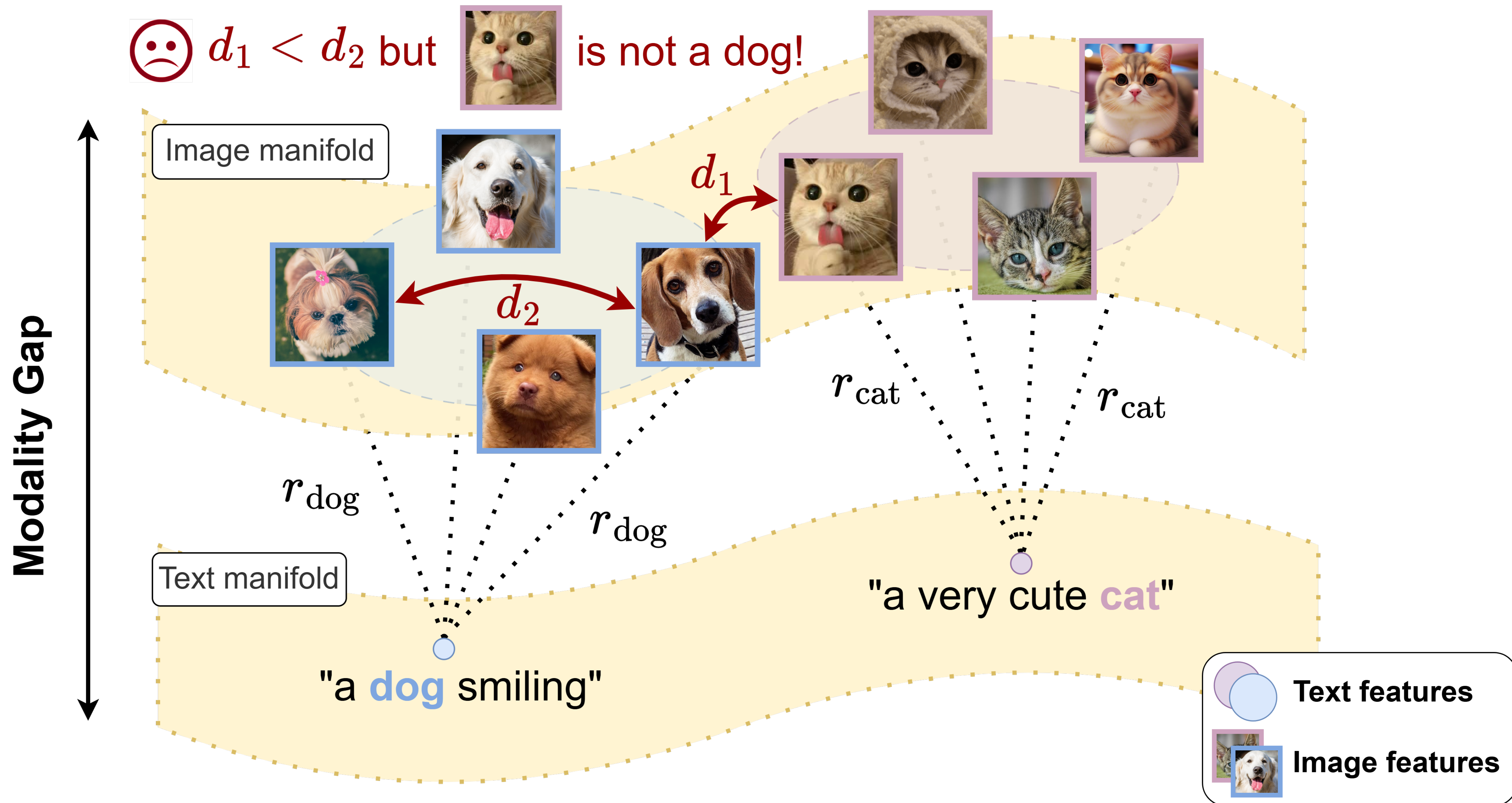
Exposing the Intra-modal misalignment

1

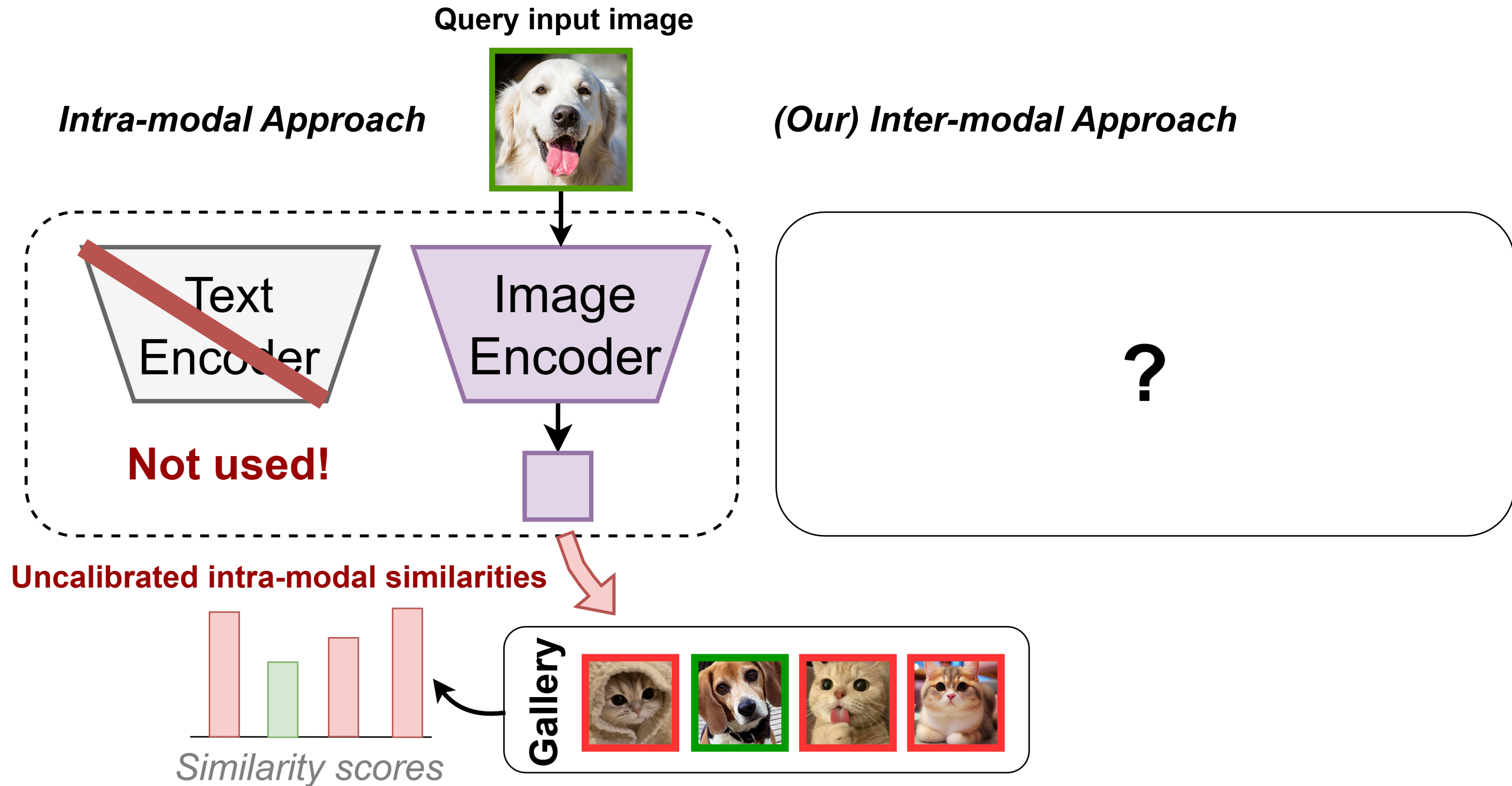


Exposing the Intra-modal misalignment

1

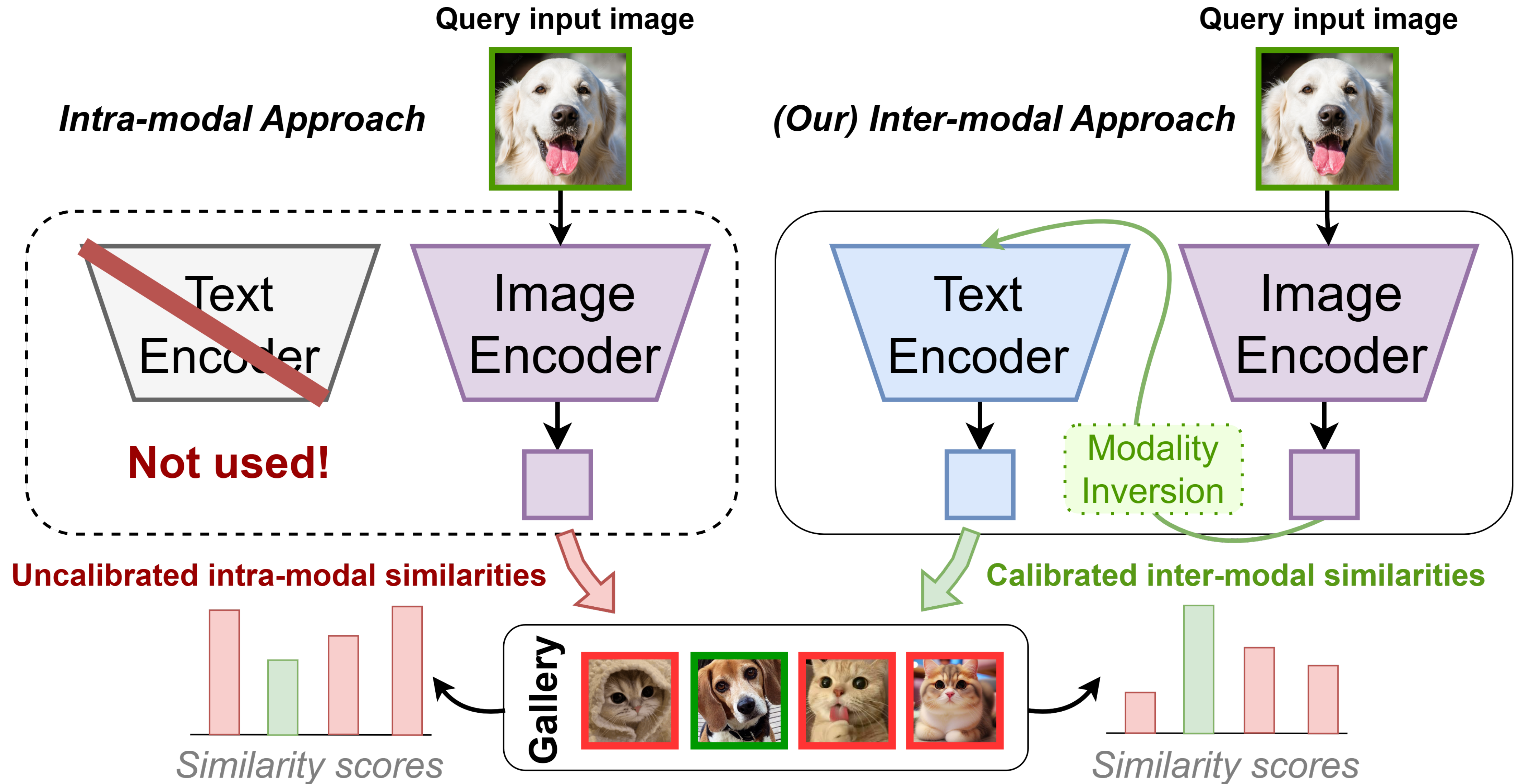


From Intra-modal to Inter-modal via Modality Inversion



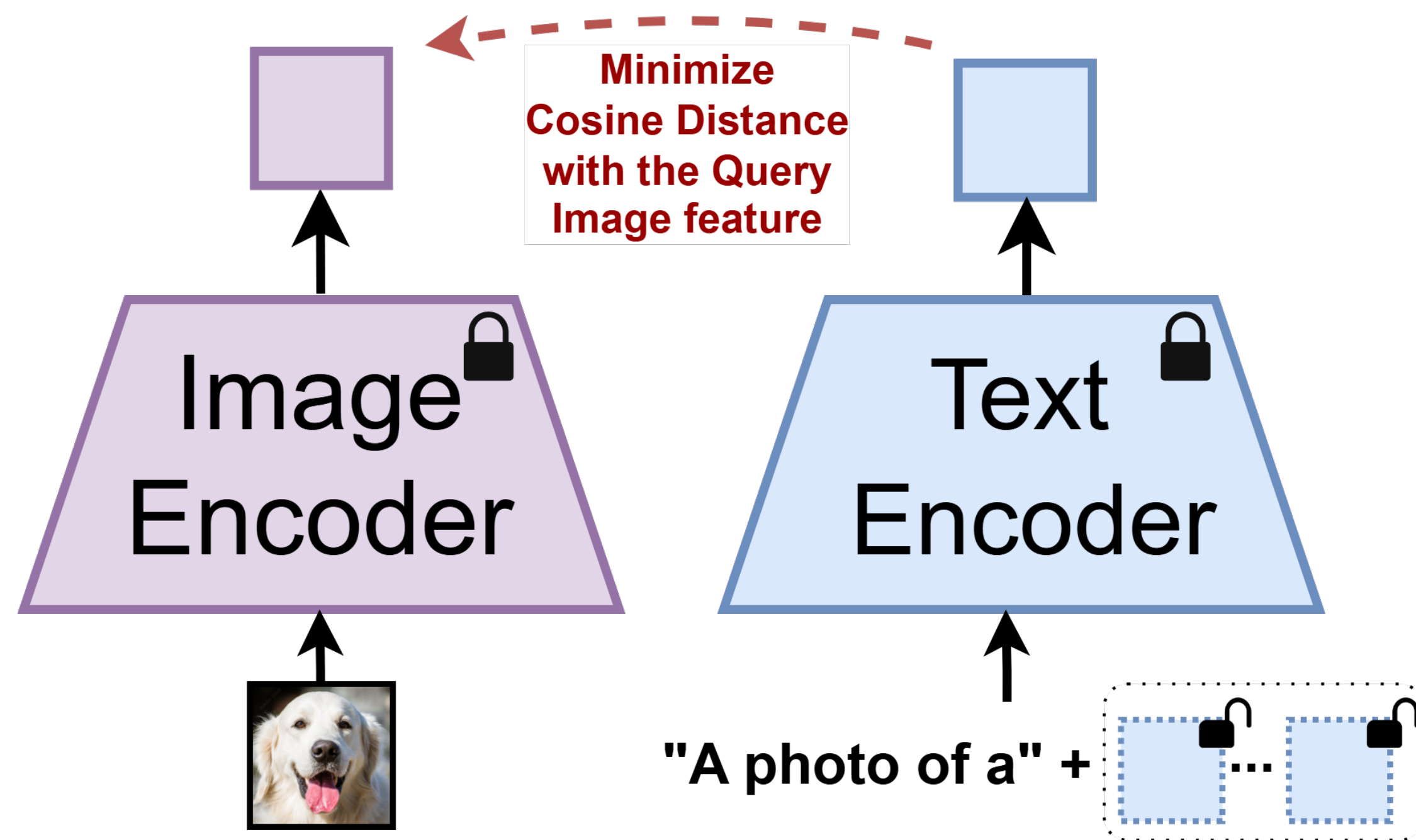
↑ Example for the intra-modal task of image-to-image retrieval. Note that text-to-text retrieval can be approached analogously

From Intra-modal to Inter-modal via Modality Inversion



↑ Example for the intra-modal task of image-to-image retrieval. Note that text-to-text retrieval can be approached analogously

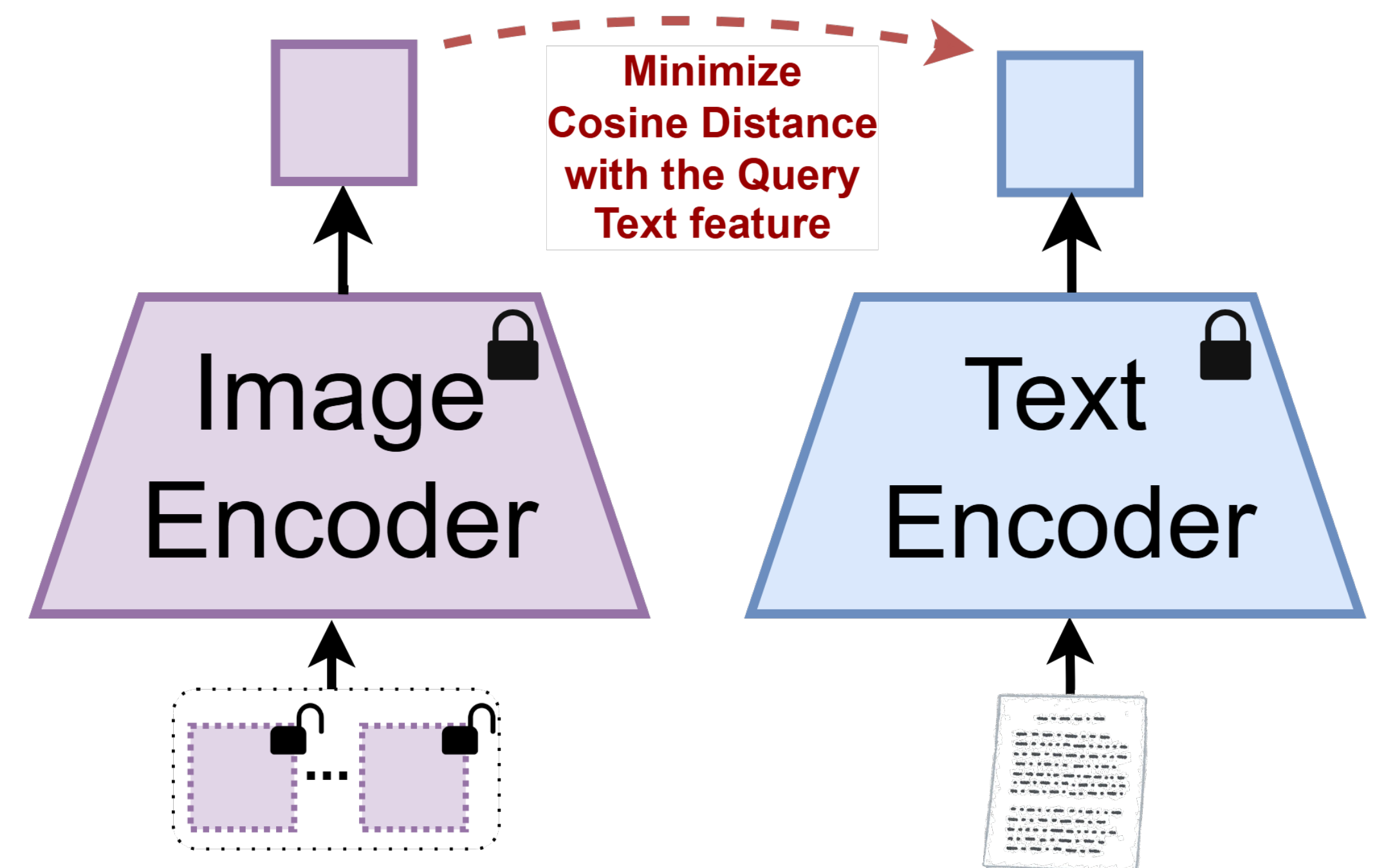
Optimization-based Textual Inversion (OTI)

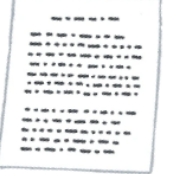


 **Pseudo token**  **Input image**

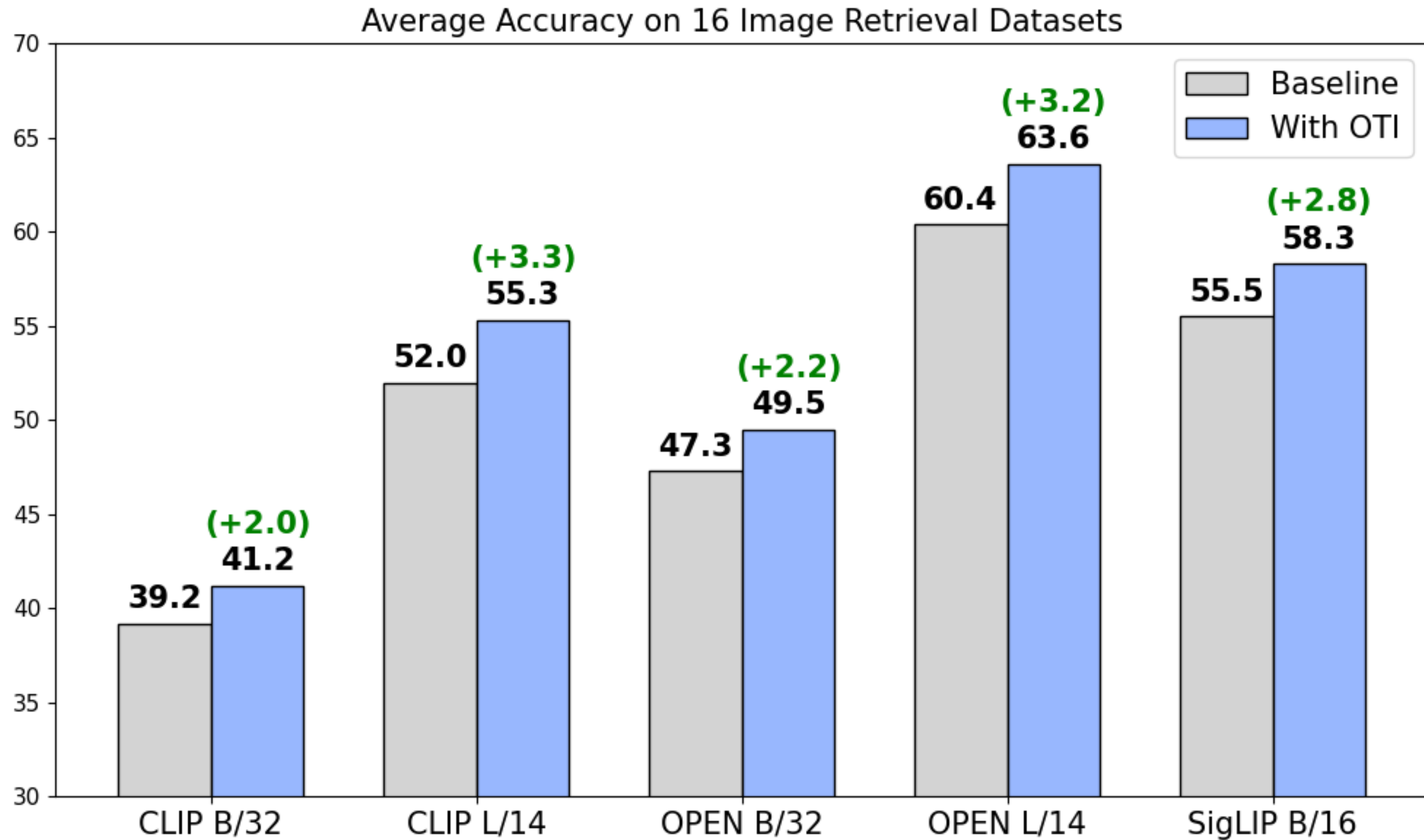
 **Frozen**  **Learnable**

Optimization-based Visual Inversion (OVI)

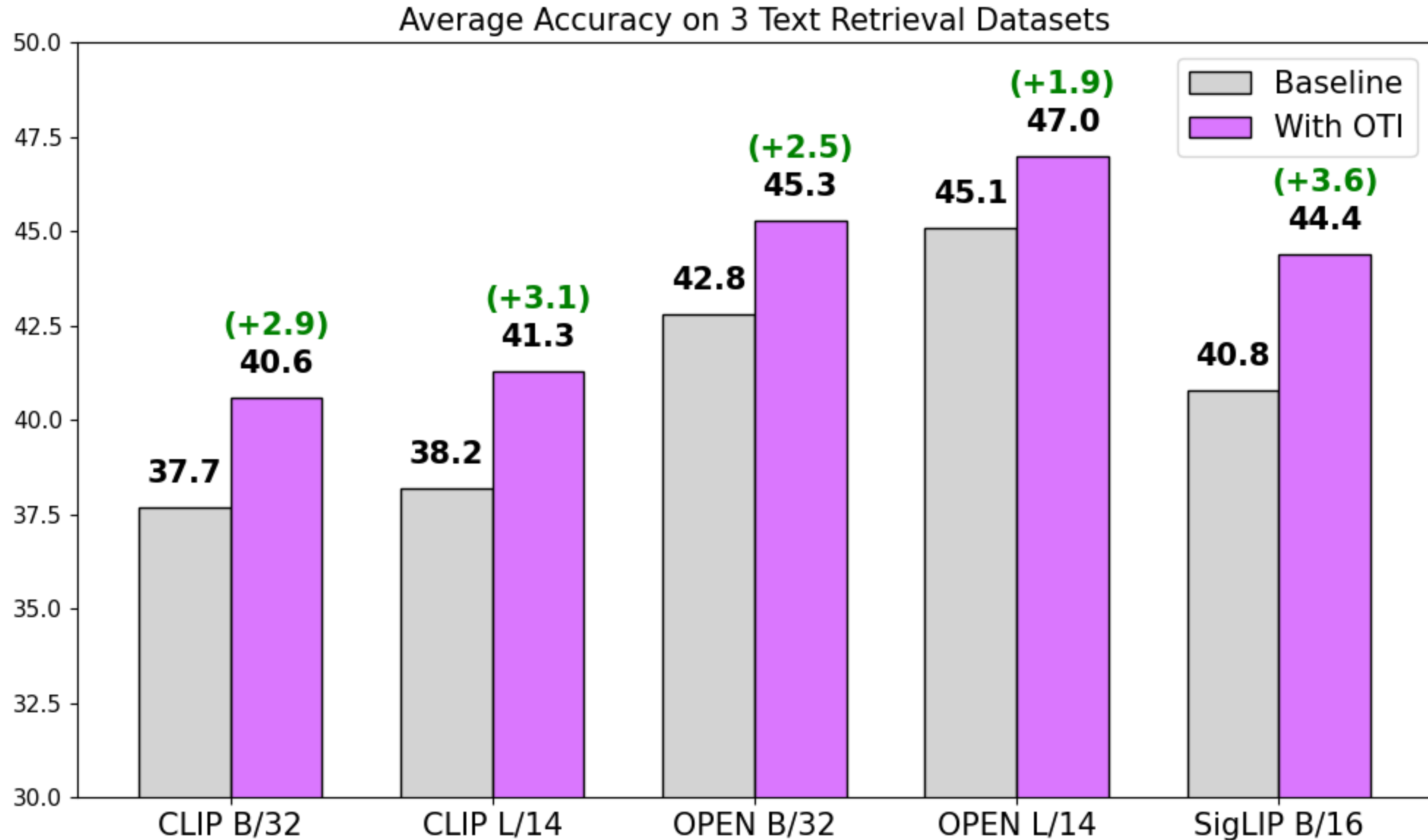


 **Pseudo patch**  **Input text**

Approaching Image-to-image retrieval inter-modally with *OTI*



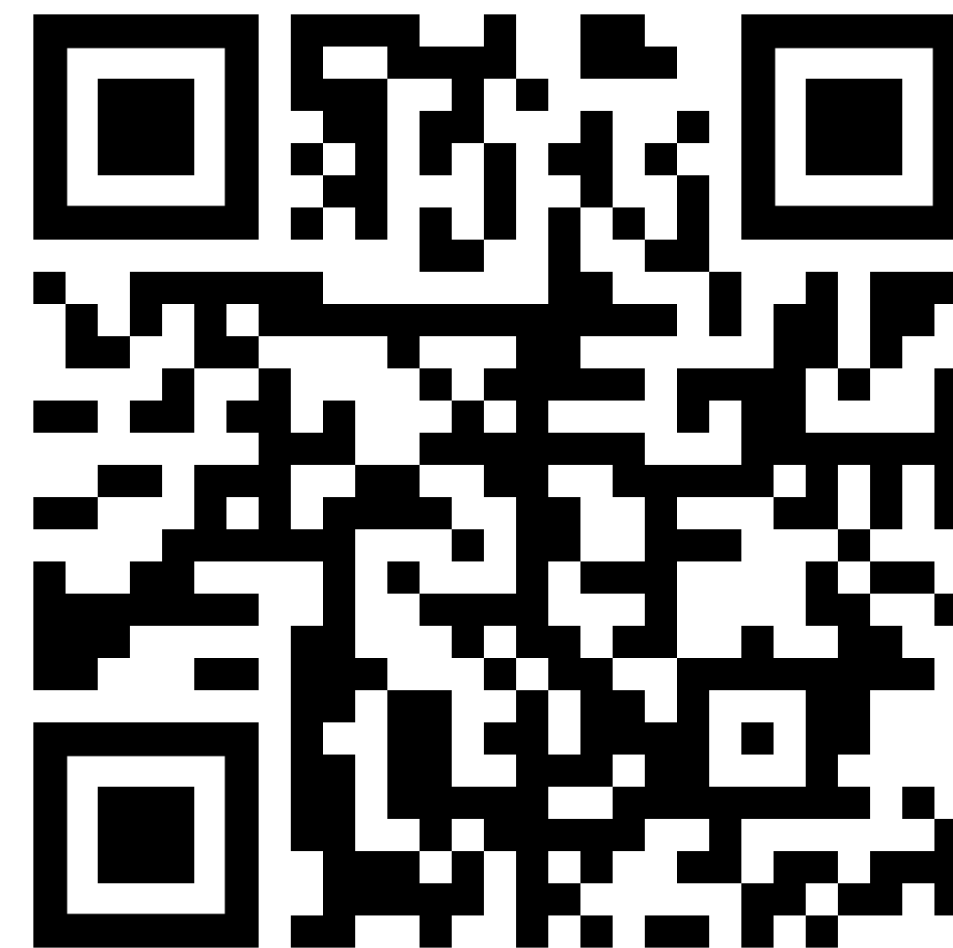
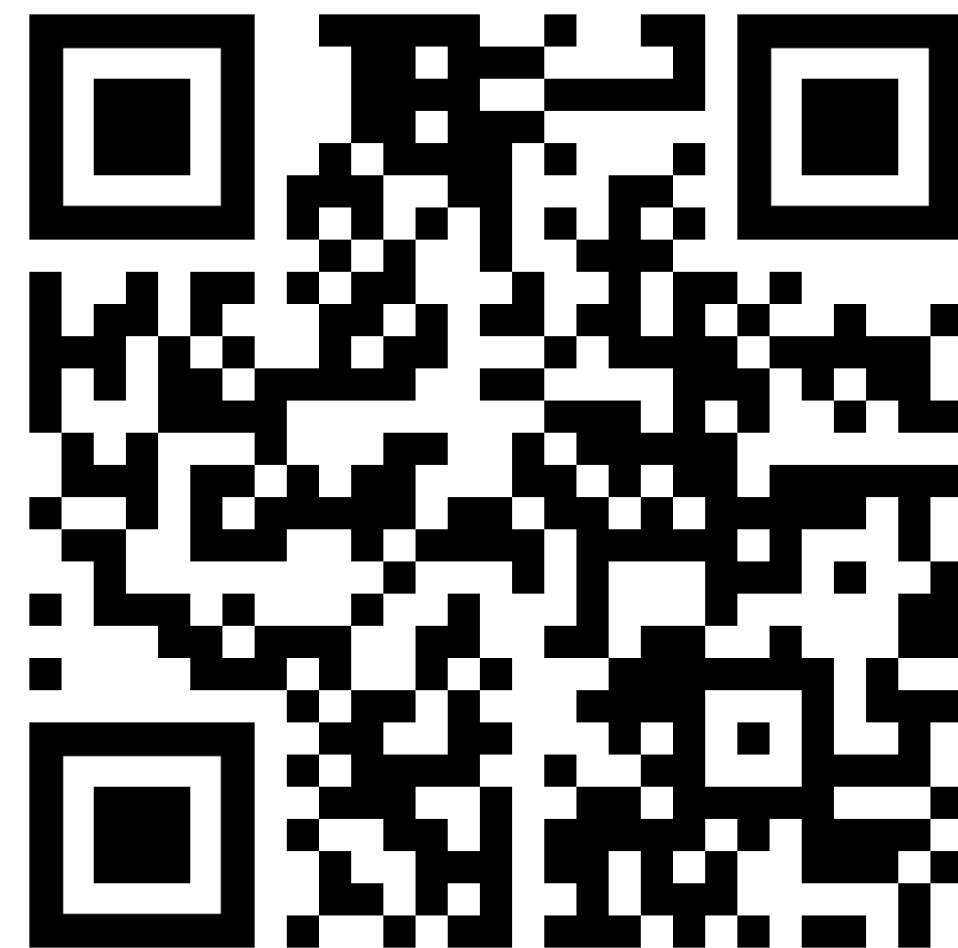
Approaching Text-to-text retrieval inter-modally with *OVI*



Cross the Gap: Exposing the Intra-modal Misalignment in CLIP via Modality Inversion

*Marco Mistretta, *Alberto Baldrati, *Lorenzo Agnolucci, Marco Bertini, Andrew D. Bagdanov

Thank you for your attention!



Presented by Marco Mistretta
marco.mistretta@unifi.it



ICLR