

APPENDIX

A TABLE OF CONTENTS

1. Sec. B provides detailed implementation information on identifying localization heads and obtaining the language-guided localization prior using these heads.
2. Sec. C describes the experimental setup, including datasets and preprocessing, baseline methods, and prompt templates.
3. Sec. D presents qualitative results of our method.
4. Sec. E discusses the limitations of our approach and directions for future work.

B IMPLEMENTATION DETAILS

In this section, we provide implementation details. We first describe how specific attention heads with strong visual grounding capabilities are selected in Sec. B.1. We then explain how text-to-image attention maps are obtained with these selected heads from different MLLMs in Sec. B.2.

B.1 FINDING LOCALIZATION HEADS

To identify localization heads (Sec. 3), we follow the approach of Kang et al. (2025), using two criteria: overall attention strength and spatial concentration.

First, for each attention head, we compute the total attention from the final text token to all image tokens and average this value over 1,000 randomly sampled queries from the RefCOCO training set (Kazemzadeh et al., 2014). Heads with average attention above the elbow point of the resulting distribution are retained. Next, to assess spatial concentration, we binarize each retained head’s attention map at the mean value, extract connected components using 8-neighbor connectivity, and compute the spatial entropy of the component size distribution (Batty, 1974; Peruzzo et al., 2022). Lower entropy indicates that attention is compactly focused on object regions, while higher entropy reflects dispersed patterns. For each query, we rank retained heads by entropy and record the top-10 with the lowest values. Finally, heads that consistently appear in these top-10 lists across queries are selected as localization heads, ensuring both strong and spatially concentrated attention.

B.2 LANGUAGE-GUIDED LOCALIZATION PRIORS

Here, we describe how the localization heads identified in the previous section is leveraged to obtain the language-guided localization prior introduced in Sec. 4.2.

By default, we adopt the top three most frequently selected heads from the previous stage, following the setting in prior work (Kang et al., 2025). Specifically, given an image and a text query, we extract the attention maps from the last text token to all image tokens of the selected localization heads. Each attention map has the shape $\mathbb{R}^{H_l W_l}$, where H_l and W_l denote the number of patches along the height and width, respectively. These maps are then reshaped to $\mathbb{R}^{H_l \times W_l}$ and further smoothed with a Gaussian filter, using a kernel size of $k = 7$ and a standard deviation of $\sigma = 1.0$. Finally, the smoothed maps are aggregated via element-wise summation to produce the final language-guided localization score map.

Among the models used in this work, Molmo-7B-D (Deitke et al., 2025) and InternVL-3-8B (Zhu et al., 2025) divide the input image into multiple square crops that tile the image, along with a global thumbnail. For Molmo-7B-D, we leverage both the attention maps to the square crops and to the thumbnail, where the thumbnail attention is interpolated to match the resolution of the square crops. For InternVL-3-8B, for simplicity, we only use the attention to the thumbnail image. In contrast, Qwen2.5-VL-7B (Bai et al., 2025), Ovis2.5-9B (Lu et al., 2025), and Kimi-VL-A3B (Team et al., 2025b) do not divide the image into square crops, so we can directly utilize the attention from text to image.

C EXPERIMENTAL DETAILS

In this section, we present the details of our experimental setup. We begin by introducing the datasets and preprocessing steps in Sec. C.1. Next, we describe the baselines used in Sec. C.2. Finally, we provide the prompt templates employed for evaluating different models in Sec. C.3.

C.1 DATASETS

In this section, we describe the datasets used in our experiments. We adapt three part segmentation datasets—PACO-LVIS (Ramanathan et al., 2023), InstructPart (Wan et al., 2025), and PartImageNet++ (Li et al., 2024)—to the few-shot part-level pointing task. For all datasets, the training split is used as the support set candidates, and the test split is used for evaluation.

PACO-LVIS contains 75 object categories and 456 part categories, covering common everyday objects. Its official test split provides 5,729 instance-level referring expressions that specify both objects and parts. For our task, we pair each referred object instance with its annotated parts to construct the evaluation set. Each sample thus consists of a target part and a referring expression for an object instance, with the goal of predicting the part of the instance. To ensure reliable evaluation, we discard samples with ground truth masks smaller than 196 pixels, resulting in 18,154 evaluation samples. Applying the same filter to the training split yields 120,777 support samples. This setup enables part-level pointing evaluation in real-world images with instance-level references.

InstructPart is designed for task-oriented part-level understanding, featuring instructions about object affordances and functionalities. It contains 2,400 images across 48 object categories and 44 part categories in household tasks, with 1,800 training (supporting) and 600 test samples. The dataset supports two tasks proposed in the original work: Task Reasoning Part Localization (e.g., “find the part in the image that can be gripped”) and Oracle Referring Part Localization (e.g., “locate the handle of the mug”). We focus on the latter, where each sample specifies an object’s part to be localized. We directly use the dataset without additional processing to evaluate models’ ability to localize parts in household scenes.

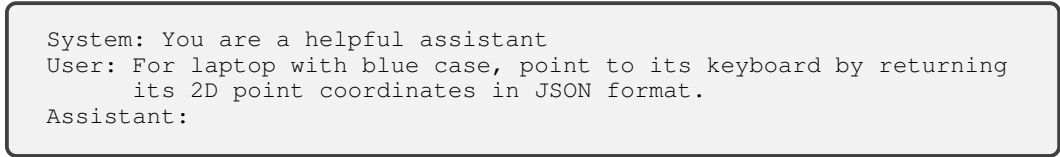
PartImageNet++ provides high-quality part segmentation annotations for all ImageNet-1K categories (Russakovsky et al., 2015), covering a diverse range of objects. The dataset defines 3,308 object-part categories, with 100 images for each of the 1,000 object categories. Each sample involves locating a specific part of an object category. We use the first 90 images per category for training (supporting) and the remaining 10 for testing, whose part annotations yield 26,747 part-level test samples. We evaluate on this dataset primarily to broaden the diversity of tested part-object categories.

C.2 BASELINES

To contextualize our few-shot results, we include several baselines. For zero-shot reference, we report the performance of pointing-capable MLLMs, including Qwen2.5-VL-7B (Bai et al., 2025), Ovis2.5-9B (Lu et al., 2025), Molmo-7B-D (Deitke et al., 2025), and the closed-source model GPT-4.1 (GPT-4.1-mini-2025-04-14).

For additional reference, we compare against several widely used open-vocabulary and reasoning segmentation models: X-Decoder (Zou et al., 2023a), SEEM (Zou et al., 2023b), VL-Part (Sun et al., 2023), and LISA (Lai et al., 2024). X-Decoder and SEEM are strong baselines for open-vocabulary and referring segmentation; we evaluate them using checkpoints with the Focal-L (Yang et al., 2022) backbone. VL-Part, in contrast, is a specialist model for open-vocabulary part segmentation, trained jointly on multiple part-level datasets (LVIS (Gupta et al., 2019), PACO, PartImageNet, and Pascal-Part) with a SwinBase Cascade Mask R-CNN architecture (Liu et al., 2021; Cai & Vasconcelos, 2018). We include it as a reference for part-focused specialists. Finally, we include LISA Lai et al. (2024), which integrates the language generation capabilities of MLLMs with mask prediction, serving as a representative grounding-specialist MLLM. LISA is also trained on part-level datasets such as PACO and Pascal-Part, and we use the LISA-7B-v1-explanatory checkpoint from the official repository.

For the few-shot baseline, we first evaluate pointing-capable MLLMs under an in-context learning setup (Brown et al., 2020), where support examples are provided directly in the context. In addition,

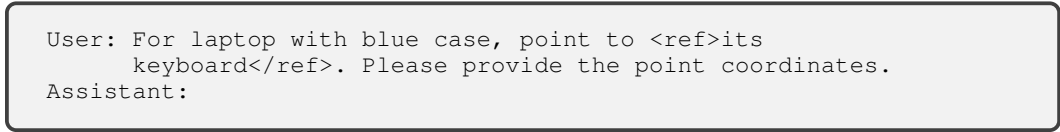


```

System: You are a helpful assistant
User: For laptop with blue case, point to its keyboard by returning
      its 2D point coordinates in JSON format.
Assistant:

```

Figure 4: Prompt format used for Qwen2.5-VL-7B evaluation on the PACO dataset.



```

User: For laptop with blue case, point to <ref>its
      keyboard</ref>. Please provide the point coordinates.
Assistant:

```

Figure 5: Prompt format used for Ovis2.5-9B evaluation on the PACO dataset.

we consider two strong training-free few-shot segmentation models, Matcher (Liu et al., 2024) and GF-SAM (Zhang et al., 2024a). Both leverage visual foundation models (DINOv2-ViT-L/14 (Oquab et al., 2024) and SAM (Kirillov et al., 2023)) to perform training-free few-shot segmentation, making them suitable baselines for our setting. For fairness, we replace DINOv2-ViT-L/14 with DINOv3-ViT-L/16 (Siméoni et al., 2025) and set the input resolution for DINOv3 to 1024 in our experiments. During evaluation, these models are provided with full support masks, whereas our method only receives a single support point, following our problem formulation.

For segmentation models that output masks, we consider two strategies for extracting a representative point: selecting the pixel with the maximum logit and selecting the innermost point determined by the distance transform (Borgefors, 1986). We observe that the maximum-logit strategy generally yields better results, so we adopt it as the default. For models that only produce binary masks, such as VL-Part, Matcher, and GF-SAM, we instead rely on the distance-transform method.

C.3 PROMPT TEMPLATE

In this section, we first describe the problem format for each dataset and then present the prompt templates used in this work. We follow the official prompt formats from the corresponding repositories or papers as closely as possible.

For evaluation, we design problem prompts tailored to each dataset. For PACO-LVIS, where referring expressions are often short sentences, we use the template: “For {referring expression}, point to its {part}.” For InstructPart and PartImageNet++, we adopt the template: “Point to the {object}’s {part}.” Only minimal modifications are made to align with the official recommendations. Across all experiments, we maintain a consistent problem format for fair comparison across models. For non-pointing-capable MLLMs, we replace “point to” with “locate” to align with their native prompt format for visual grounding tasks.

For each model, we follow its official prompt specification as closely as possible. Using the PACO dataset and the query “For the laptop, point to its keyboard.” as an example, Fig. 4, Fig. 5, and Fig. 6 show the exact prompt formats used to evaluate Qwen2.5-VL-7B, Ovis2.5-9B, and Molmo-7B-D respectively. For non-pointing-capable MLLMs, Fig. 7 and Fig. 8 present the corresponding formats used for InternVL-3-8B and Kimi-VL-A3B.

D QUALITATIVE RESULTS

We present several qualitative results to illustrate the effectiveness of our approach. Fig. 9, Fig. 10, and Fig. 11 show representative examples from Molmo-7B-D, Ovis2.5-9B, and Qwen2.5-VL-7B on PACO dataset. Ground-truth regions are highlighted with red masks, and for clearer visualization, the prediction is shown on a cropped image.

User: For laptop with blue case, point to its keyboard.
Assistant:

Figure 6: Prompt format used for Molmo-7B-D evaluation on the PACO dataset.

User: For laptop with blue case, locate the region this sentence
describes: <ref>its keyboard</ref>.
Assistant:

Figure 7: Prompt format used for InternVL-3-8B evaluation on the PACO dataset.

E LIMITATIONS AND FUTURE WORK

In this section, we discuss the limitations of our approach and potential directions for future work.

1. Our current evaluation is limited to tasks requiring a single-point prediction. Extending the method to multi-target scenarios—for example, through non-maximum suppression (NMS)—is an important avenue for future work.
2. The effectiveness of our approach depends on the quality of the underlying attention maps. When they are not well aligned with semantic parts, incorporating few-shot exemplars yields only modest gains.
3. Predicted points correspond to patch centers, so localization accuracy is constrained by encoder resolution: even when the correct patch is identified, the predicted center may deviate from the true location.

User: For laptop with blue case, locate its keyboard.
Assistant:

Figure 8: Prompt format used for Kimi-VL-A3B evaluation on the PACO dataset.

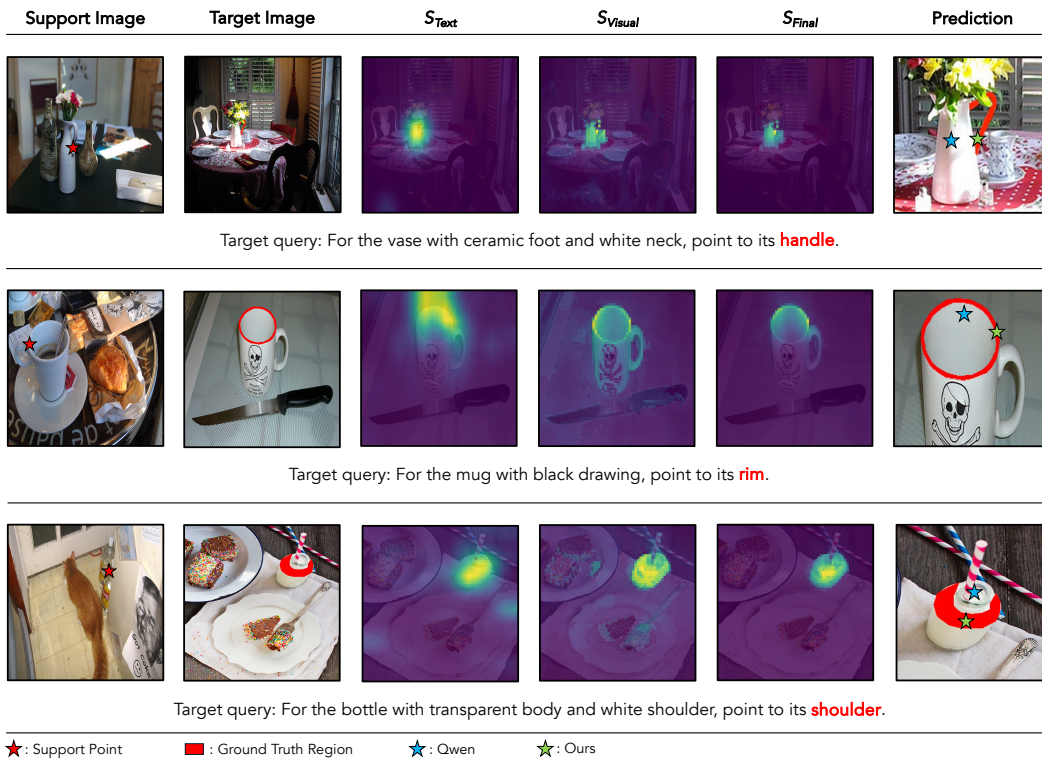


Figure 9: Qualitative results of Qwen2.5-VL-7B on PACO. For clearer visualization, the prediction is shown on a cropped image.

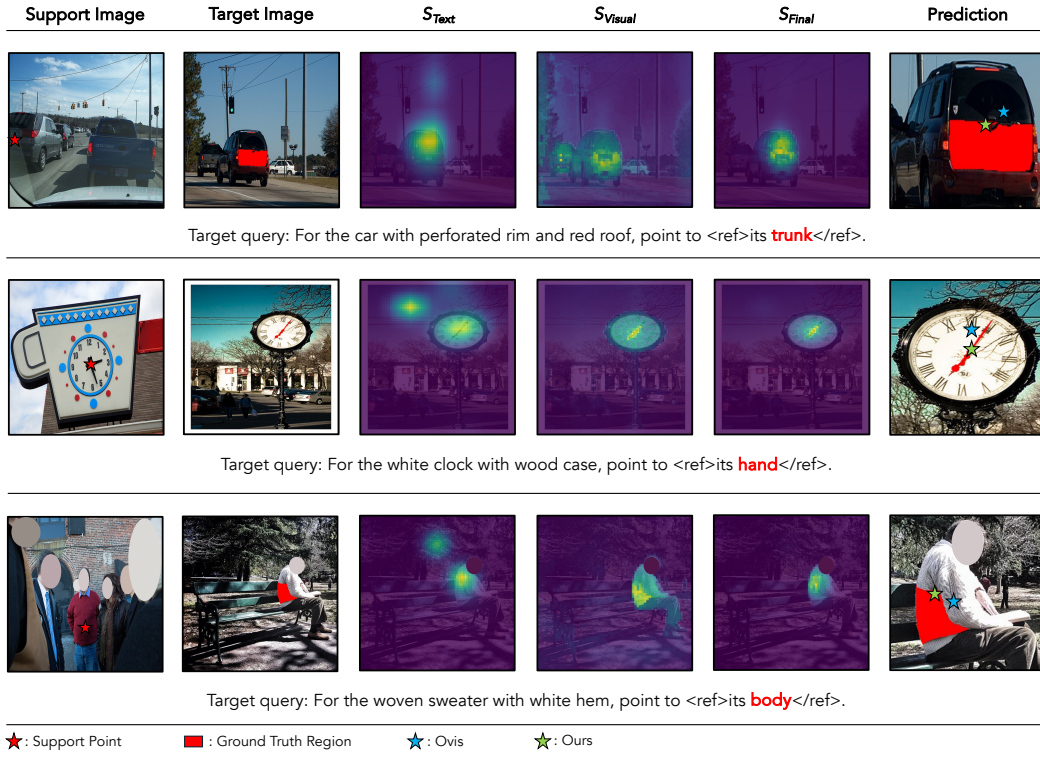


Figure 10: Qualitative results of Ovis2.5-9B on PACO. For clearer visualization, the prediction is shown on a cropped image.

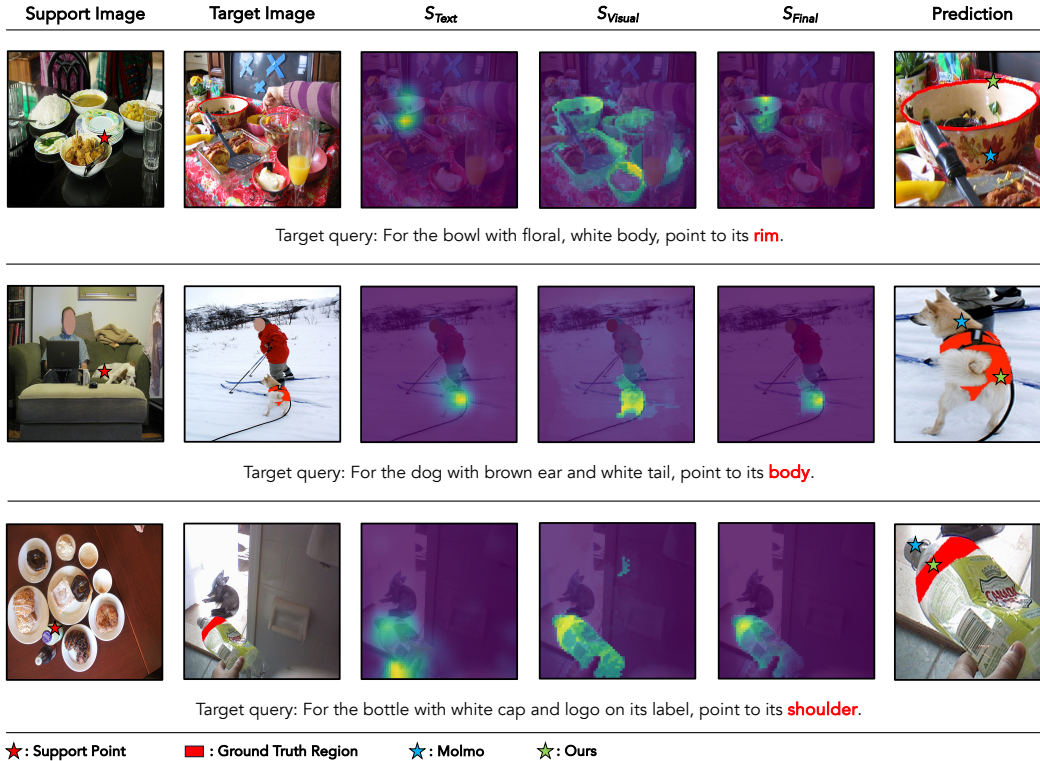


Figure 11: Qualitative results of Molmo-7B-D on PACO. For clearer visualization, the prediction is shown on a cropped image.