

InstructSAM: A Training-Free Framework for Instruction-Oriented Remote Sensing Object Recognition

Yijie Zheng^{1,2} Weijie Wu^{1,2} Qingyun Li³ Xuehui Wang⁴ Xu Zhou⁵
Aiai Ren⁵ Jun Shen⁵ Long Zhao¹ Guoqing Li¹ Xue Yang⁴
¹Aerospace Information Research Institute ²University of Chinese Academy of Sciences
³Harbin Institute of Technology ⁴Shanghai Jiao Tong University ⁵University of Wollongong



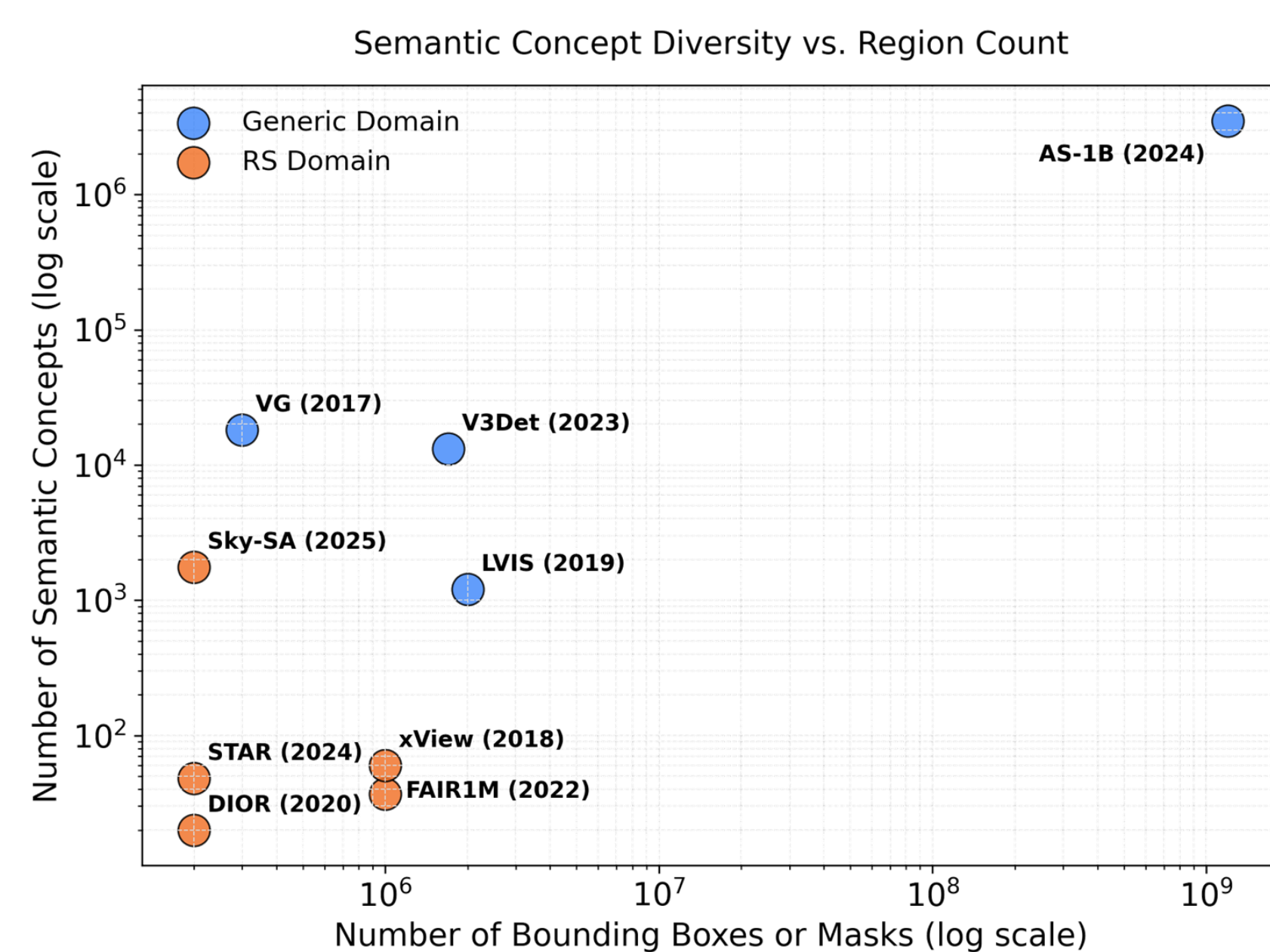
From close-set to instruction-oriented

- Fixed categories limit traditional detection.
- Real tasks need flexible, dynamic recognition.
- ✧ **Instruction-oriented: just give an instruction, model adapts.**



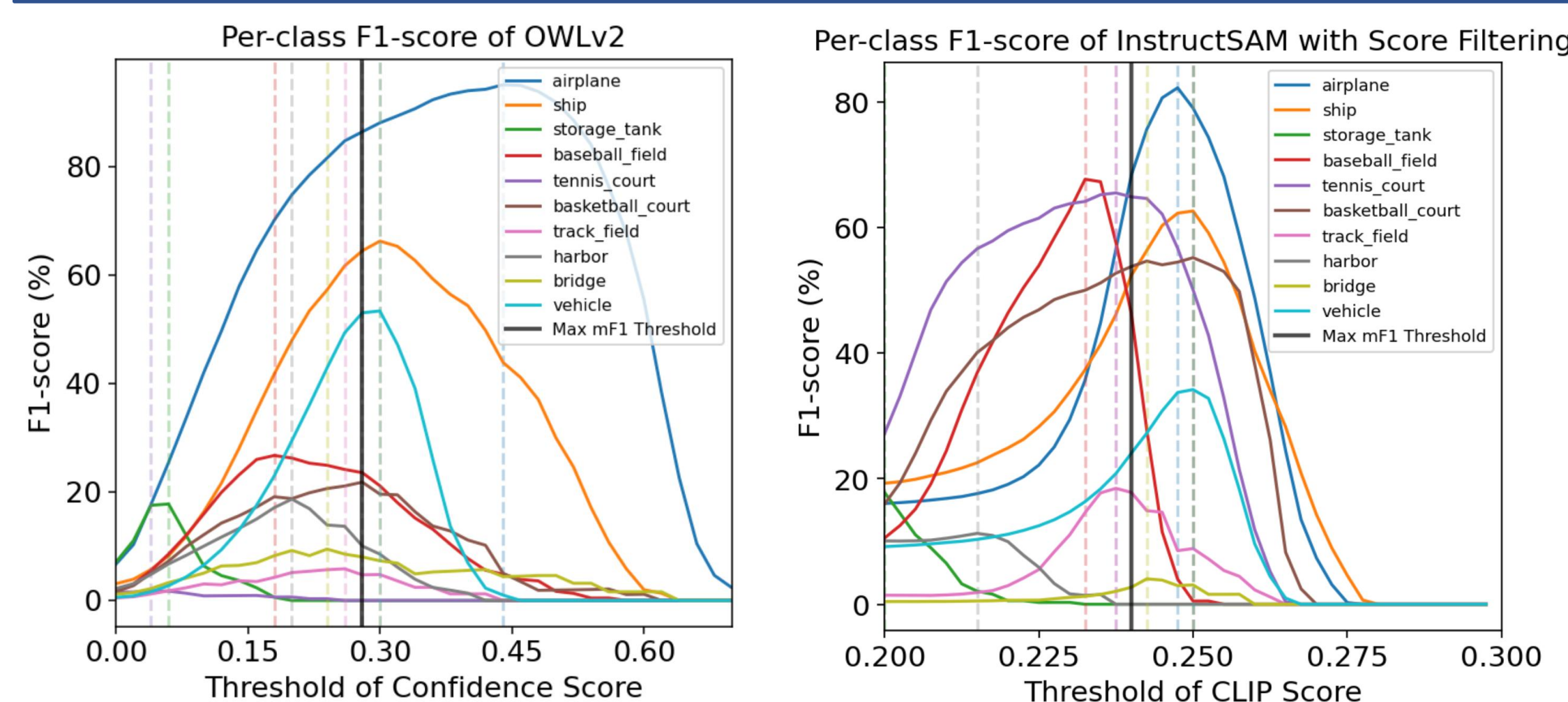
Close-Set	Open-Vocabulary	Open-Ended	Open-Subclass	...
R-CNN CVPR'14	OVR-CNN CVPR'21	GenerateU CVPR'24	Ins-DetCLIP ICLR'24	InstructSAM NeurIPS'25

Motivation 1: Lacking diverse training datasets



- RS datasets lack semantic diversity → poor generalization.
- 💡 **Use generic LVLMs to identify objects.**

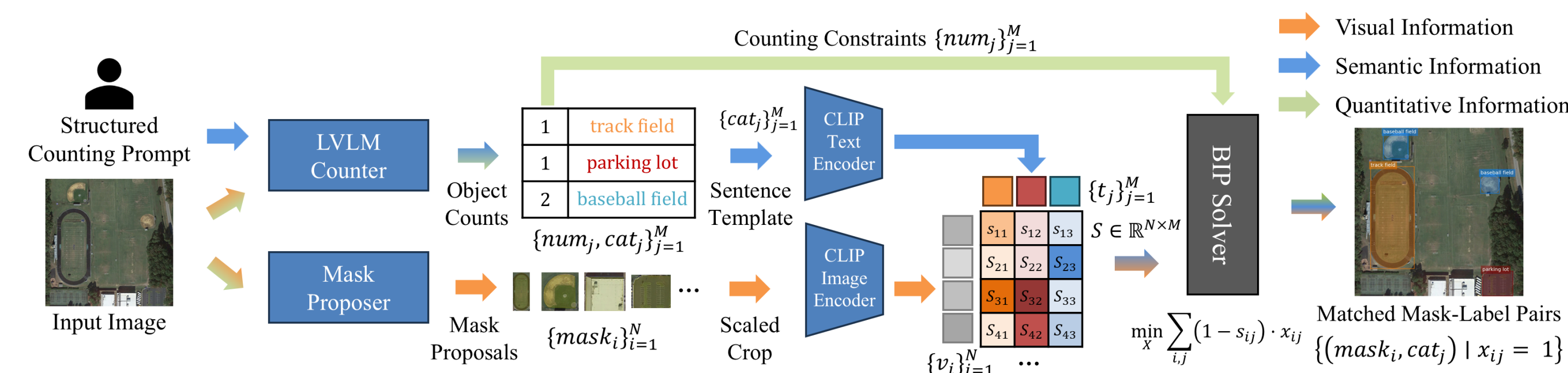
Motivation 2: Unreliable score filtering



- Score-based filtering is crucial to filter low-quality pseudo labels.
- Optimal thresholds vary across classes → no universal solution.
- Over-reliance on score thresholds leads to misclassifications.
- 💡 **Introduce counting constraints**

InstructSAM framework

- Decompose object segmentation into three easier steps.
- 💡 LVLm for categories & counts, SAM for mask proposals, CLIP for similarity, PuLP for optimization.



- Reframe segmentation as a mask-label matching problem

$$\min_{\mathbf{x}} \sum_{i=1}^N \sum_{j=1}^M (1 - s_{ij}) \cdot x_{ij} \quad \bullet \text{ minimize mismatches (1 - similarity)}$$

$$\text{s.t.} \quad \sum_{j=1}^M x_{ij} \leq 1, \quad \bullet \text{ One mask} \rightarrow \text{one category}$$

$$\sum_{i=1}^N x_{ij} = num_j, \quad \bullet \text{ Total masks per category} = \text{LVLM count}$$

LVLM counts precisely and quickly

- With clear rules, GPT-4o counts as accurately as Faster R-CNN (80% vs. 81% mF1).



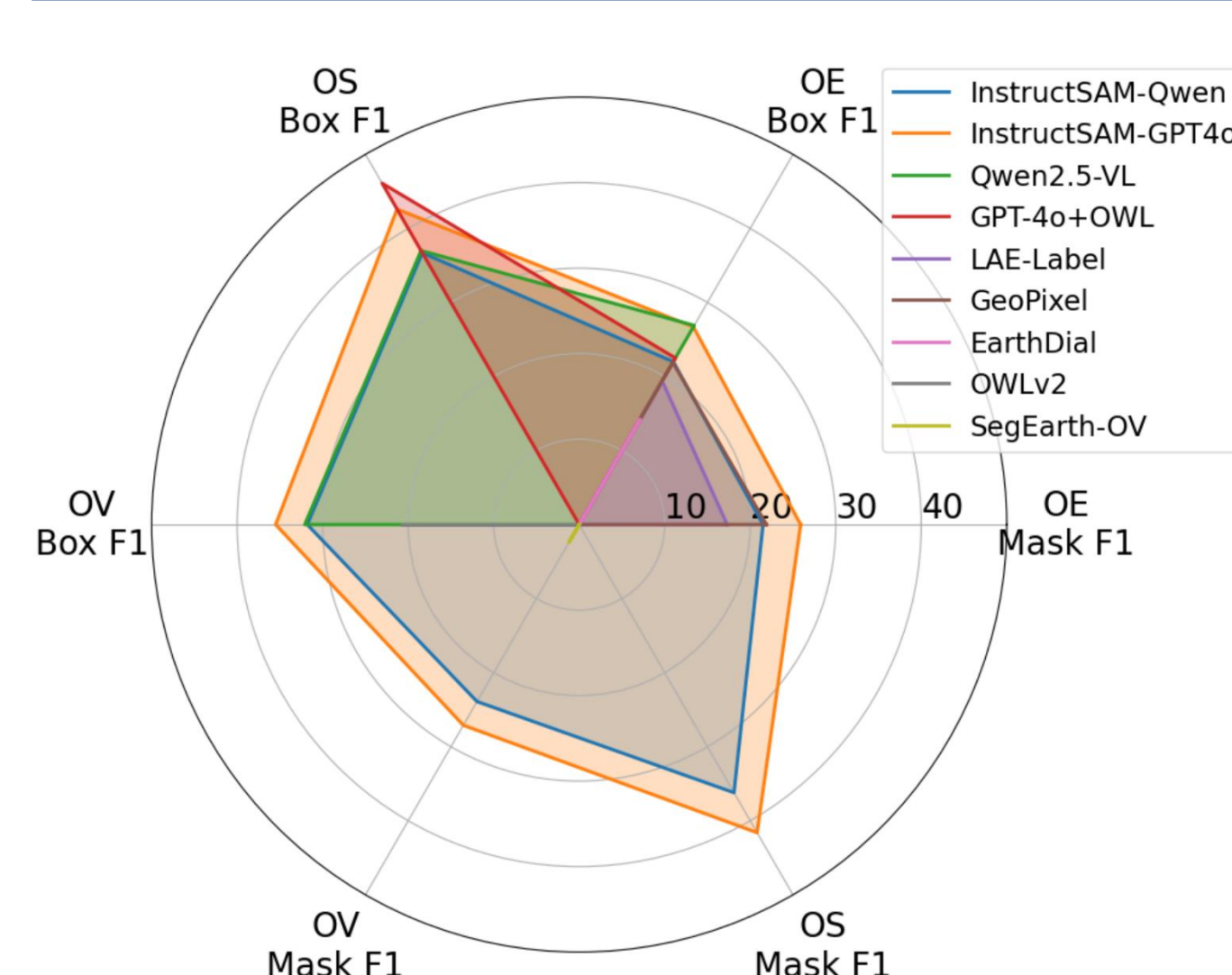
- Count the number of harbors. Answer in JSON format.
- ✗ {"harbor": 1}
- Count the number of harbors. Answer in JSON format.
- Instructions:
- Harbor = pier to dock ships, Count each pier separately.
- ✓ {"harbor": 8}

- Counting greatly reduces tokens and inference time.

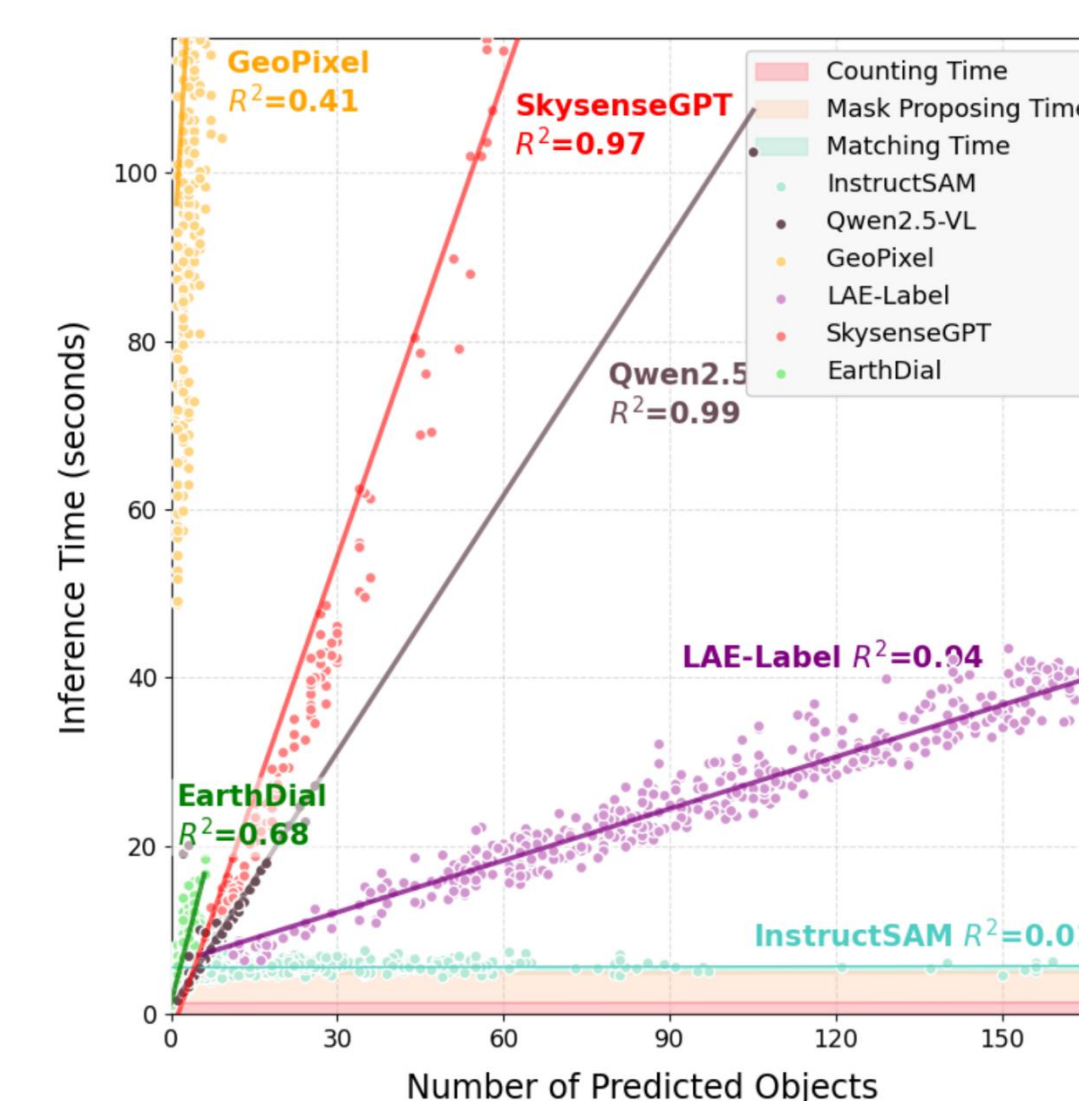


- Count the number of planes.
- ✓ {"plane": 8} (6 tokens, 0.5s)
- Detect all the planes.
- [{"bbox_2d": [58, 40, 78, 97], "label": "plane"}, {"bbox_2d": [52, 56, 70, 60], "label": "plane"}, ... {"bbox_2d": [42, 22, 76, 33], "label": "plane"}] (183 tokens, 10s)

Zero-Shot Results across Three Settings



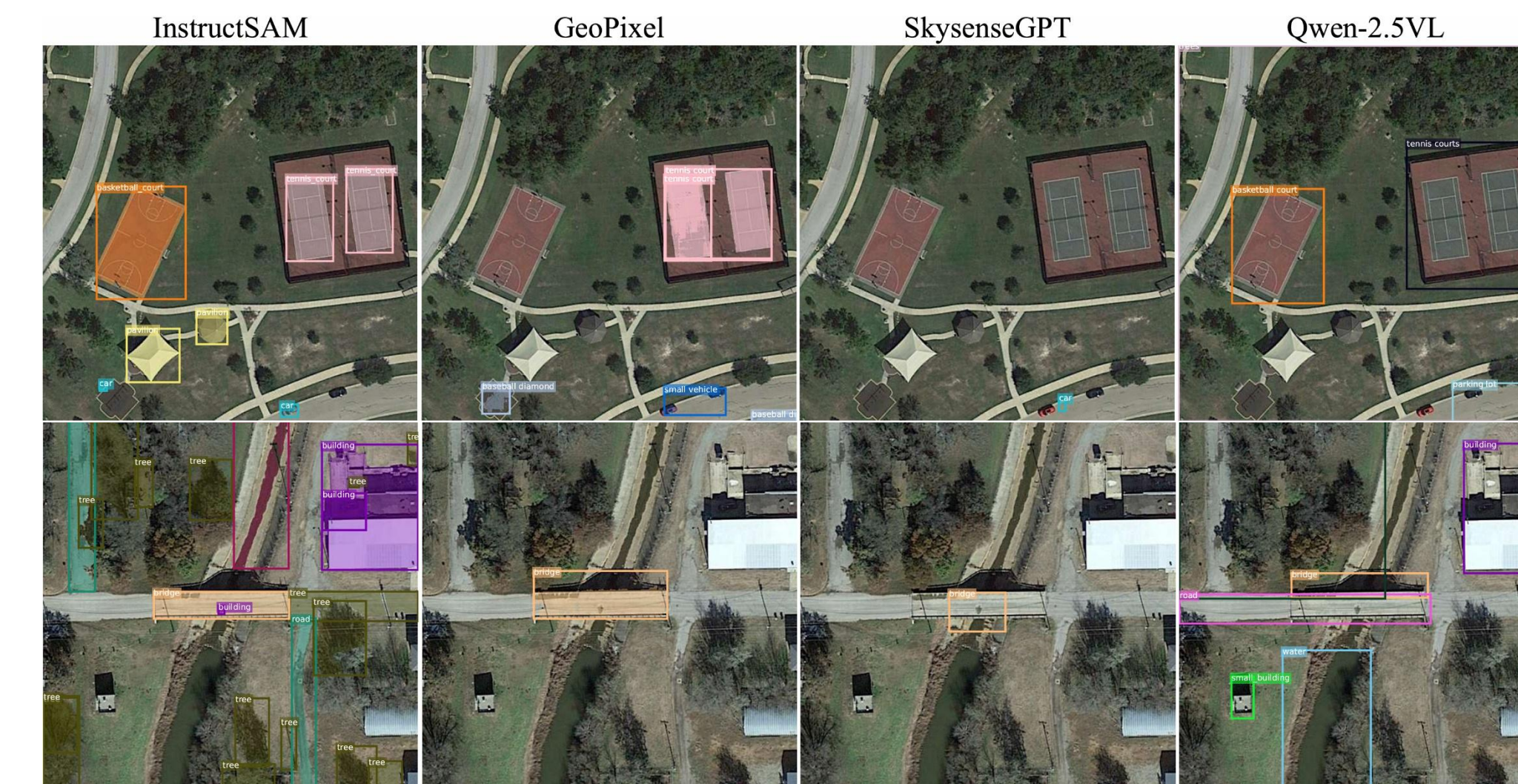
- Evaluated on NWPU and DIOR dataset
- 🔥 Strong performance across most settings
- 👉 Performs well using open models (Qwen2.5-VL)



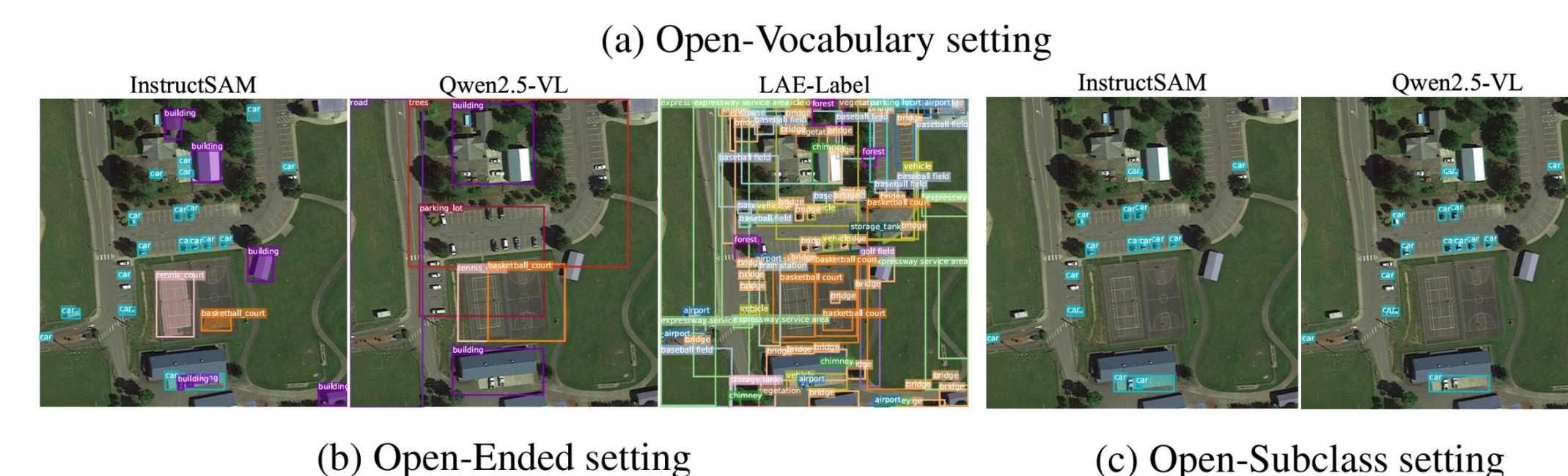
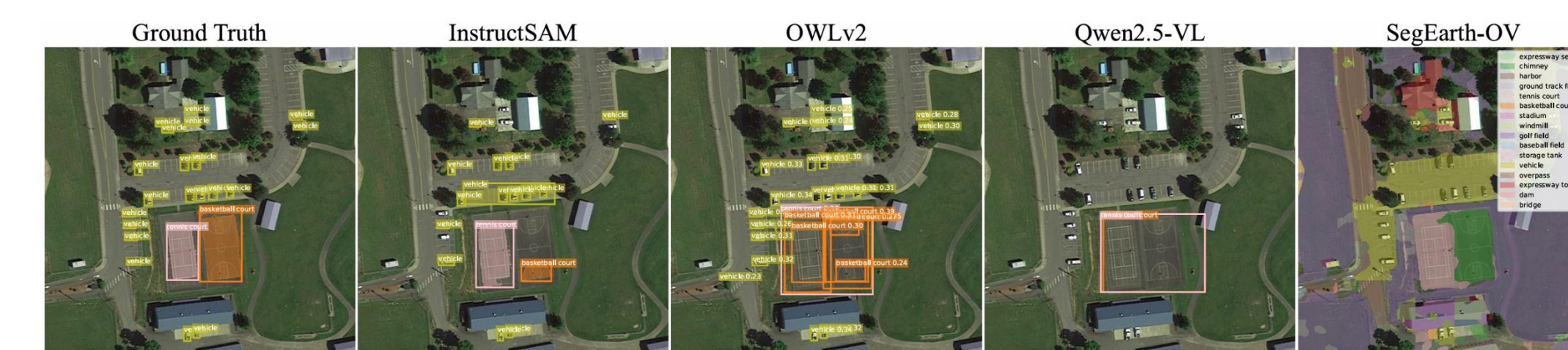
- Inference time in open-ended setting
- ⚡ 89% fewer tokens, 32% inference time reduction
- 🕒 Inference time stays constant with more objects

Qualitative results: open-ended setting

- ✗ Existing RSVLMs fail to generalize beyond training categories.
- ✓ InstructSAM recognizes more categories, e.g., pavilion and tree.



Qualitative results across three settings

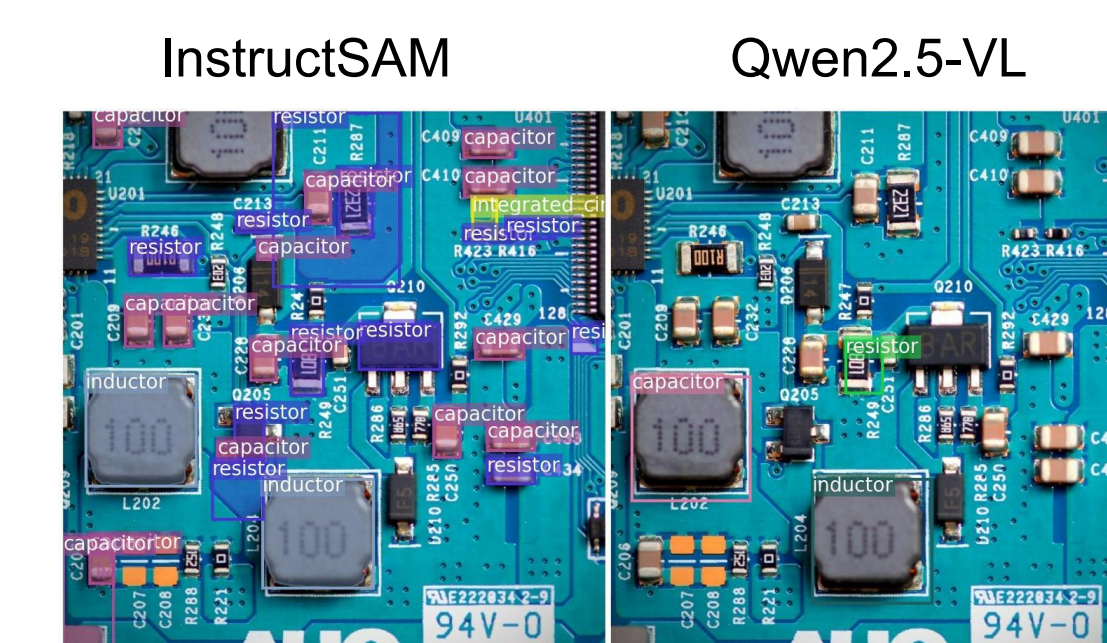


Generalization beyond RS

- InstructSAM also generalizes to natural images.

Instruction:

Detect all electronic components.



Instruction:

Detect dices whose letters come before K.



Takeaways & future work

Takeaways

- Flexible: Works with diverse instructions.
- Efficient: Faster than detection, saves tokens, scales well.
- Training-Free: Directly benefits from stronger open or proprietary models.

Future Work

- Expand InstructSAM to semantic segmentation (land cover mapping).
- Develop stronger Remote Sensing Foundation Models (e.g., SAM, CLIP).



Code



Personal Page

- Open for PhD / Visiting Opportunities
- zhengyijie23@mails.ucas.ac.cn