

Recurrent Cross-View Object Geo-Localization

Xiaohan Zhang¹, Si-Yuan Cao¹, Xiaokai Bai¹, Yiming Li¹, Zhangkai Shen¹, Zhe Wu¹, Xiaoxi Hu², Hui-liang Shen¹

¹College of Information Science and Electronic Engineering, Zhejiang University

²State Key Laboratory of Intelligent Green Vehicle and Mobility, Tsinghua University, Tsinghua University
zhangxh2023@zju.edu.cn

Abstract

Cross-view object geo-localization (CVOGL) aims to determine the location of a specific object in high-resolution satellite imagery given a query image with a point prompt. Existing approaches treat CVOGL as a one-shot detection task, directly regressing object locations from cross-view information aggregation, but they are vulnerable to feature noise and lack mechanisms for error correction. In this paper, we propose **ReCOT**, a **Recurrent Cross-view Object geo-localization Transformer**, which reformulates CVOGL as a recurrent localization task. ReCOT introduces a set of learnable tokens that encode task-specific intent from the query image and prompt embeddings, and iteratively attend to the reference features to refine the predicted location. To enhance this recurrent process, we incorporate two complementary modules: (1) a SAM-based knowledge distillation strategy that transfers segmentation priors from the Segment Anything Model (SAM) to provide clearer semantic guidance without additional inference cost, and (2) a Reference Feature Enhancement Module (RFEM) that introduces a hierarchical attention to emphasize object-relevant regions in the reference features. Extensive experiments on standard CVOGL benchmarks demonstrate that ReCOT achieves state-of-the-art (SOTA) performance while reducing parameters by 60% compared to previous SOTA approaches.

1 Introduction

Cross-view object geo-localization (CVOGL) aims to determine the geographic location of a specific object indicated by point prompts in a query image on the reference image (Sun et al. 2023). The query images can be captured from devices like phones, autonomous vehicles, robots, and drones, while the reference images are typically high-resolution satellite images. CVOGL is widely used in various applications, such as smart city management (Yao et al. 2022), disaster monitoring (Chini, Pierdicca, and Emery 2009; Kumar, Kim, and Hancke 2013), and robot navigation (Singamaneni et al. 2024; Zhai et al. 2024). However, the view gap poses challenges for CVOGL (Sun et al. 2023).

Recent cross-view image geo-localization (CVIGL) works (Hu et al. 2018; Shi et al. 2019; Zhu, Shah, and Chen 2022; Yang, Lu, and Zhu 2021; Lin et al. 2022) have demonstrated their superiority in handling view gaps. However,

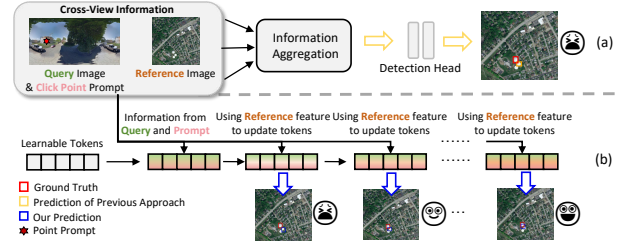


Figure 1: Comparison between the framework of previous CVOGL approaches and ours. (a) Previous approaches treat the CVOGL as a prompt-based detection task, where the model directly regresses the object location based on information aggregation once. (b) Our framework reformulates the CVOGL as a recurrent localization problem, where the model iteratively refines the localization through a set of learnable tokens. Please refer to the zoomed-in view for better visualization.

CVIGL approaches are fundamentally designed for camera-level localization using retrieval-based approaches (Deuser, Habel, and Oswald 2023; Zhang et al. 2024b; Shi et al. 2019) or fine-grained approaches (Sarlin et al. 2023; Wang et al. 2023). However, CVOGL aims to localize specific objects (e.g., a building with a red roof) captured in the query image, which demands prompt-guided and object-aware prediction. Therefore, in CVOGL scenarios, CVIGL approaches can only provide a nearby location for the indicated object (Sun et al. 2023), which is insufficient for precise object-level localization.

To address this, CVOGL approaches emerge recently. Existing approaches (Sun et al. 2023; Li et al. 2025; Huang et al. 2025) typically treat CVOGL as a one-shot detection paradigm, where the model directly regresses the object location based on prompt-guided information aggregation, as shown in Fig. 1(a). For example, the recent state-of-the-art (SOTA) approach (Huang et al. 2025) aggregates information from cross-view images and prompts to produce a spatial attention matrix, which is used to enhance the reference image features. The enhanced features are then fed into several convolutional layers to regress the object location. While efficient and architecturally simple, such a framework is sensitive to the quality of the enhanced feature (Cao et al.

2022). It lacks a correction mechanism for early-stage prediction errors, making them vulnerable to noise in features, *i.e.*, the model cannot give correct localization once the enhanced feature leads to a wrong prediction (Cao et al. 2023; Sun et al. 2023).

To cope with this, we propose a **Recurrent Cross-view Object geo-localization Transformer (ReCOT)**. Motivated by the success of iterative refinement strategies in matching tasks (Yu et al. 2023; Cao et al. 2022, 2023), ReCOT reformulates the CVOGL task as a recurrent localization problem, as shown in Fig. 1(b). This serves as the main difference between previous approaches (Sun et al. 2023; Li et al. 2025; Huang et al. 2025) and our framework. Specifically, ReCOT initializes a set of learnable tokens, which interact with the query image feature and prompt embeddings to extract task-specific intent. These tokens then act as recurrent “questioners” that iteratively attend to the enhanced reference image features, progressively extracting object-relevant information and refining the prediction. Both the interaction and refinement processes are implemented using a combination of self- and cross-attention mechanisms. The proposed recurrent strategy enables our ReCOT to effectively enhance the performance of CVOGL by iteratively refining the initial prediction, as shown in Fig. 1(b).

Nevertheless, the success of this recurrent strategy in our token-driven framework relies on the semantic clarity of the prompts (Li et al. 2025) and the quality of reference features (Teed and Deng 2020). Specifically, the learnable tokens are first guided by prompt semantics to extract object-relevant intent, then iteratively attend to the reference features to refine the object location in each recurrent step. Therefore, if prompts lack clear semantic intent or reference features are cluttered, the tokens may not accumulate correct task-specific cues across iterations, leading to suboptimal performance. To tackle this, we introduce two complementary methods: (1) a SAM-based knowledge distillation strategy, which injects prior knowledge from large-scale model into the prompt embeddings to boost prompt understanding while avoiding computational cost during inference, and (2) a Reference Feature Enhancement Module (RFEM), which emphasizes object-relevant reference features through hierarchical attention. These components provide clean visual and semantic cues, enabling the tokens to effectively accumulate task-specific information during iterative refinement.

We evaluate ReCOT on the standard CVOGL benchmark (Sun et al. 2023). It achieves state-of-the-art (SOTA) performance while reducing parameter count by 60% compared to the previous SOTA approach (Huang et al. 2025) (29.9M vs. 74.8M), and runs at a competitive inference speed. In summary, our contributions are as follows:

- We propose ReCOT, a novel framework for CVOGL that reformulates the task as a recurrent localization problem, where learnable tokens iteratively attend to reference features to refine object localization.
- We introduce a SAM-based knowledge distillation strategy that transfers prior knowledge from a large foundation model into the prompt embeddings, providing clearer semantic guidance without adding inference cost.

- We design the RFEM, which leverages a proposed hierarchical attention to highlight object-relevant regions in the reference feature, thereby facilitating the recurrent localization process.

2 Related Work

Cross-View Image Geo-Localization (CVIGL). CVIGL aims to determine the camera’s geographic location by matching a ground-view query image with the most correlated reference image (Deuser, Habel, and Oswald 2023; Zhang et al. 2024b; Shi et al. 2019) or position (Sarlin et al. 2023; Wang et al. 2023; Lentsch et al. 2023). Existing CVIGL approaches can be grouped into metric learning-based methods (Lu, Luo, and Zhu 2022; Zhu et al. 2022; Cai et al. 2019; Shi et al. 2020b; Guo et al. 2022; Shi and Li 2022; Shi et al. 2022; Hu et al. 2018; Yang, Lu, and Zhu 2021; Lin et al. 2022; Zhu, Shah, and Chen 2022), which learn viewpoint-invariant features, and geometry-based methods (Shi et al. 2020a; Lu et al. 2020; Toker et al. 2021; Liu and Li 2019; Regmi and Shah 2019; Shi et al. 2019), which exploit orientation or structural cues to reduce viewpoint gaps. However, CVIGL methods only provide camera-level localization and cannot accurately pinpoint object-level targets (Sun et al. 2023).

Cross-View Object Geo-Localization (CVOGL). CVOGL focuses on locating a specific object indicated by prompts in a query image. DetGeo (Sun et al. 2023) first formalized this task and proposed a detection-based framework. Subsequent works, such as VaGeo (Li et al. 2025) and OCGNet (Huang et al. 2025), enhanced cross-view feature aggregation and prompt embedding. Despite progress, existing CVOGL methods still rely on one-shot detection, which is sensitive to noisy features and lacks mechanisms for error correction (Cao et al. 2022). In contrast, we reformulate CVOGL as a recurrent localization problem and propose ReCOT to address these limitations.

3 Methodology

Fig. 2 presents the architecture of ReCOT, which reformulates CVOGL as a recurrent localization task. A set of learnable tokens is initialized to encode task-specific intent from the query image features and prompt embeddings. These tokens act as recurrent “questioners” that iteratively extract object-relevant cues from the reference features and refine the predicted location through cross-attention mechanisms. Additionally, we introduce two complementary methods to enhance this recurrent process: (1) a SAM-based knowledge distillation strategy, which transfers segmentation priors from the Segment Anything Model (SAM) to enhance prompt semantics, and (2) a Reference Feature Enhancement Module (RFEM), which provides object-relevant reference features through proposed hierarchical attention for the recurrent stage.

Recurrent Localization Framework

Motivation. As shown in Fig. 1, existing CVOGL approaches follow a one-shot detection paradigm, directly predicting the object location from enhanced reference features.

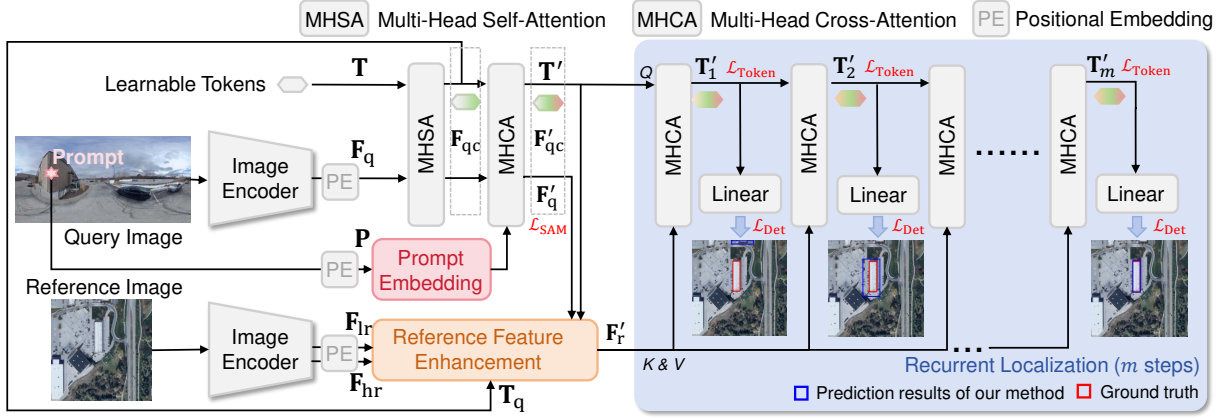


Figure 2: Overall architecture of our recurrent cross-view object geo-localization transformer (ReCOT). ReCOT reformulates the CVOGL task as a recurrent localization task, which leverages a set of learnable tokens to extract information from cross-view images and prompts to recurrently refine the prediction. Notably, all recurrent steps in ReCOT share the same “MHCA” and “Linear” component.

However, such frameworks are often sensitive to feature noise and lack a mechanism for error correction (Cao et al. 2022, 2023). Fundamentally, CVOGL can be regarded as a cross-view matching problem, where recurrent strategies have shown superior robustness across domains (Cao et al. 2022; Teed and Deng 2020). Inspired by this, we reformulate CVOGL as a recurrent localization process. Moreover, unlike dense matching tasks (Cao et al. 2023; Yu et al. 2023; Edstedt et al. 2024), CVOGL is prompt-driven and focuses on object-level semantic matching. This calls for a representation that can both encode semantic intent and drive iterative refinement. To this end, we draw inspiration from the class token in vision transformers (ViT) (Dosovitskiy et al. 2021), and introduce a set of learnable tokens that absorb task-specific semantics from the query and prompt. Acting as semantic carriers, these tokens can recurrently interact with the reference feature to enable step-wise prediction.

Structure. We initialize a set of learnable tokens $T \in \mathbb{R}^{n \times c}$, where n and c denote the number of tokens and the feature dimension, respectively. To enable T to extract object-relevant information from the reference features, T needs to first acquire task-specific semantics from the query image feature $F_q \in \mathbb{R}^{h_q w_q \times c}$ and the point prompt embedding $P \in \mathbb{R}^c$. Here, h_q and w_q denote the height and width of the query feature, respectively. Specifically, we concatenate T with F_q along the spatial dimension, yielding $F_{qc} \in \mathbb{R}^{(n+h_q w_q) \times c}$. Following standard operations in Vision Transformers (ViT) (Dosovitskiy et al. 2021), we apply self-attention to F_{qc} , allowing T to aggregate global context from the query image. The resulting tokens are denoted as T_q . To further inject object-level intent, the point prompt P embedded interacts with F_{qc} through cross-attention. This enables the prompt to semantically guide the tokens, allowing T_q to further incorporate object-level context. After interaction, we denote the concatenated feature and tokens as F'_{qc} and T' , respectively. The F'_{qc} and T_q are further utilized to enhance reference features in RFEM. The enhanced ref-

erence feature are denoted as $F'_r \in \mathbb{R}^{h_r w_r \times c}$, where h_r and w_r denote the height and width of the reference feature, respectively

After acquiring task-specific semantics from the query and prompt, we use T' to perform recurrent localization, as shown in Fig. 2. Let T'_i denote the token state at the i -th refinement step, where $i \in [0, 1, 2, \dots, m]$. We set m to 6 in our work experimentally. At each step, T'_i attends to the enhanced reference feature F'_r to extract object-relevant cues and update the task-specific intent. Formally, the update process is defined as

$$T'_{i+1} = \text{MHCA}(T'_i, F'_r) \quad (1)$$

where $\text{MHCA}(\cdot, \cdot)$ denotes the multi-head cross-attention module. Here, T'_i serves as the query, while F'_r acts as the key and value. This recurrent attention mechanism enables iterative refinement, where the tokens T' progressively accumulate task-specific semantics and extract increasingly relevant information from the fixed reference feature F'_r . The T'_i of each step is fed into a linear layer to predict an updated object location, allowing the model to gradually converge toward a precise localization.

At each refinement step i , we introduce a loss $\mathcal{L}_{Token}$ to guide the generation of T'_i . Specifically, since T'_i is expected to contain the task-specific intent in cross-view features, it should be able to highlight the required object area on F'_r . Therefore, in each refinement step, we first aggregate T'_i along the spatial dimension to generate a global embedding T''_i , and use $T''_i \in \mathbb{R}^{1 \times c}$ and F'_r to produce an aggregation map $\hat{m}_o \in \mathbb{R}^{1 \times h_r \times w_r}$. This can be expressed as

$$T''_i = \text{Sum}(T'_i \cdot \text{Softmax}(T'_i)), \quad (2)$$

$$\hat{m}_{oi} = \sigma(T''_i F_r^T), \quad (3)$$

where $\sigma(\cdot)$ and $\text{Softmax}(\cdot)$ denote the sigmoid and softmax function, respectively. $\text{Sum}(\cdot)$ is the summing along the spatial dimension. We then utilize a box-level mask m_{oi} produced using the ground truth box to supervise the generation

of $\hat{\mathbf{m}}_o$, which can be expressed as

$$\mathcal{L}_{\text{Token}_i}(\mathbf{m}_o, \hat{\mathbf{m}}_{oi}) = \mathcal{L}_{\text{bce}}(\mathbf{m}_o, \hat{\mathbf{m}}_{oi}) + \mathcal{L}_{\text{dice}}(\mathbf{m}_o, \hat{\mathbf{m}}_{oi}), \quad (4)$$

where $\mathcal{L}_{\text{bce}}(\cdot, \cdot)$ and $\mathcal{L}_{\text{dice}}(\cdot, \cdot)$ denote the binary cross-entropy loss and the Dice Loss (Milletari, Navab, and Ahmadi 2016), respectively.

How ReCOT works. As shown in Fig. 3(a), previous one-shot detection CVOGL approaches (Sun et al. 2023; Li et al. 2025; Huang et al. 2025) rely on a single forward information aggregation and are thus sensitive to noisy features (Cao et al. 2022, 2023). Our ReCOT adopts a recurrent localization mechanism that iteratively refines predictions. The visualization in Fig. 3(b) of cross-attention weights between tokens and the reference feature reveals the inner dynamics of this refinement process. It can be seen that different tokens focus on different regions of the reference feature, indicating a form of token-level specialization. For object-relevant tokens, their attention gradually concentrates and intensifies around the object region across recurrent steps, reflecting the ability to correct early prediction error and progressively refine the prediction. In contrast, object-irrelevant tokens experience a decrease in their attention responses and eventually stabilize to background patterns once they no longer contribute to the object localization. This behavior highlights the competitive nature of token updates, *i.e.*, multiple tokens initially compete to explain different parts of the reference feature (Carion et al. 2020), but those correlated with the object receive positive feedback and their attention weights are amplified over recurrent steps, leading to iterative convergence. Such dynamics demonstrate the effectiveness of recurrent refinement mechanism in suppressing irrelevant regions while enhancing object-relevant cues.

SAM-based Knowledge Distillation

Motivation. Point prompt understanding is essential for CVOGL to correctly locate objects. However, point prompt itself suffers from semantic ambiguity, leading to unsatisfactory performance (Kirillov et al. 2023). To address this, we propose a SAM-based knowledge distillation strategy to boost the prompt understanding of ReCOT. The incorporation of SAM (Kirillov et al. 2023) is motivated by an observation that SAM can give a mask with a clear indication of the required object using point prompts and corresponding images. However, directly applying SAM during inference incurs a large computation overhead. Hence, we leverage predictions of SAM as supervision signals, transferring its knowledge through knowledge distillation.

Structure. We extract the query feature $\mathbf{F}'_q$ from the previous concatenated feature $\mathbf{F}'_{qc}$. Notably, the $\mathbf{F}'_q$ has been interacted with prompt embedding and is supposed to contain the object-level semantic. Therefore, we process it using a lightweight convolutional head followed by a sigmoid activation to generate a segmentation map $\hat{\mathbf{m}}_q$. Meanwhile, we use SAM to generate a pseudo ground-truth mask $\mathbf{m}_{\text{SAM}}$ from the query image and point prompt. This mask is used to supervise the segmentation out of $\mathbf{F}'_q$ via $\mathcal{L}_{\text{SAM}}$, which can be defined as

$$\mathcal{L}_{\text{SAM}}(\mathbf{m}_{\text{SAM}}, \hat{\mathbf{m}}_q) = \mathcal{L}_{\text{bce}}(\mathbf{m}_{\text{SAM}}, \hat{\mathbf{m}}_q) + \mathcal{L}_{\text{dice}}(\mathbf{m}_{\text{SAM}}, \hat{\mathbf{m}}_q), \quad (5)$$

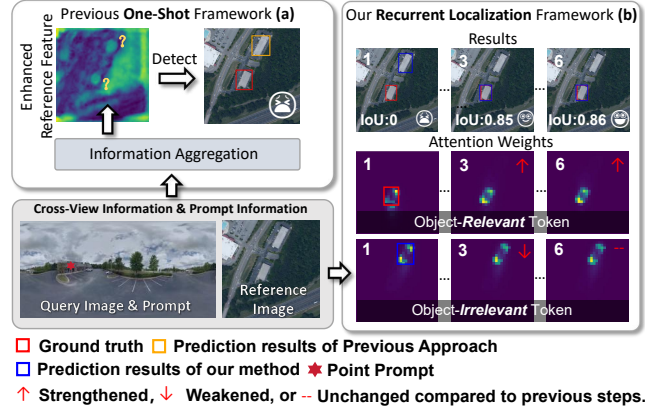


Figure 3: Comparison between the previous one-shot framework and our recurrent localization framework. (a) Previous approaches (Li et al. 2025; Sun et al. 2023; Huang et al. 2025) rely on single-shot information aggregation, which is sensitive to noisy enhanced features and often leads to incorrect predictions. (b) Our ReCOT iteratively refines predictions through learnable tokens. The attention weight visualizations show that object-relevant tokens progressively focus and strengthen around the target, while object-irrelevant tokens weaken and stabilize to background patterns. Please refer to the zoomed-in view for better visualization.

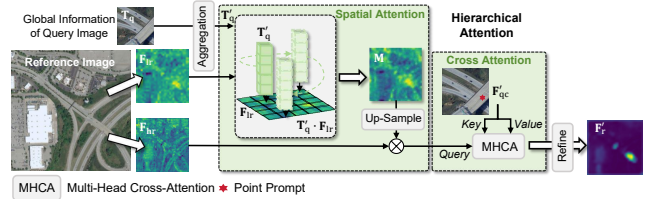


Figure 4: Architecture of reference feature enhancement module (RFEM). RFEM enhances reference features through a hierarchical attention pipeline to obtain $\mathbf{F}'_r$.

Reference Feature Enhancement Module

Motivation. In CVOGL, the reference feature $\mathbf{F}_r$ extracted by the backbone encoder is typically generic and background-dominated, lacking object-level specificity before prompt interaction (Sun et al. 2023; Li et al. 2025). In our framework, guiding $\mathbf{F}_r$ to focus on the expected object indicated by the prompt can significantly ease the downstream recurrent localization (Cao et al. 2023; Teed and Deng 2020). To this end, we propose the Reference Feature Enhancement Module (RFEM), which enhances reference features into a more object-aware representation $\mathbf{F}'_r$. Unlike previous approaches (Sun et al. 2023; Li et al. 2025; Huang et al. 2025) that attempt to clearly highlight the object features through one-shot feature enhancement, RFEM serves as a preparatory module to filter irrelevant features and provide more object-relevant information for subsequent recurrent localization. The key of RFEM lies in its hierarchical attention design. We first perform spatial attention to highlight the query-relevant regions in the reference feature, as

the expected object is likely confined to these regions. We then apply cross-attention guided by the query and prompt cues to refine these regions with object-level semantics. This spatial-to-cross hierarchy yields a cleaner and more informative reference feature $\mathbf{F}'_r$, which significantly benefits the iterative localization process of ReCOT.

Structure. Specifically, as illustrated in Fig. 4, we utilize two levels of reference features, *i.e.* a low-resolution semantic feature $\mathbf{F}_{lr} \in \mathbb{R}^{h_r \times w_r \times c}$ that captures scene semantics, and a high-resolution detailed feature $\mathbf{F}_{hr} \in \mathbb{R}^{2h_r \times 2w_r \times c}$ that preserves fine-grained spatial structures (Zhang et al. 2024c).

Spatial Attention: We leverage the spatial attention to highlight query-relevant features. Specifically, We first aggregate the query tokens $\mathbf{T}_q$ into a global descriptor $\mathbf{T}'_q$ (Eq. (2)). Notably, the $\mathbf{T}'_q$ contains only the global information of the query image without prompt guidance. This descriptor correlates with $\mathbf{F}_{lr}$ via a dot-product operation to produce a spatial attention map as

$$\mathbf{M} = \sigma(\mathbf{T}'_q \mathbf{F}_{lr}^T), \quad (6)$$

which highlights query-relevant regions. The attention map is then up-sampled and applied to $\mathbf{F}_{hr}$ by element-wise multiplication, suppressing irrelevant background and narrowing the focus to potential object areas.

Cross Attention: We leverage the cross-attention to incorporate detailed query features and prompt semantics for further refining the reference feature. The filtered $\mathbf{F}_{hr}$ is subsequently refined via cross-attention with $\mathbf{F}_{qc}$, which encodes fine-grained prompt cues and detailed query features. This cross-attention step sharpens the object-relevant responses and injects object-level semantics into the reference representation.

Finally, the updated feature is down-sampled and refined through spatial attention (Woo et al. 2018) and self-attention (Dosovitskiy et al. 2021) to generate the $\mathbf{F}'_r$, which is then passed to the recurrent localization stage in ReCOT.

Why multi-resolution reference features? The reference image typically contains both global scene information and fine-grained object details. The low-resolution semantic reference feature $\mathbf{F}_{lr}$ extracted from deeper layers of the backbone encodes more scene-level context but lose spatial details due to down-sampling. Therefore, we utilize $\mathbf{F}_{lr}$ to enhance query-relevant regions. Conversely, the high-resolution reference features $\mathbf{F}_{hr}$ preserve fine structural details but are dominated by background noise. Thus, it needs to be refined by the attention map produced by $\mathbf{F}_{lr}$, and then RFEM can utilize it to perform object-level enhancement. This design leverages advantages of multi-resolution features, which is crucial for object-level feature enhancement in CVOGL (Huang et al. 2025).

Loss Function

Our training objective $\mathcal{L}$ is defined as a weighted sum of the localization loss $\mathcal{L}_{\text{local}}$ and the auxiliary loss $\mathcal{L}_{\text{aux}}$. It can be defined as

$$\mathcal{L} = \mathcal{L}_{\text{local}} + \alpha \mathcal{L}_{\text{aux}}, \quad (7)$$

where $\mathcal{L}_{\text{local}}$ denotes the sum of DETR-style detection losses $\mathcal{L}_{\text{Det}}$ computed at each recurrent localization step (see

Steps m	Ground→Satellite		Drone→Satellite	
	Acc@0.25	Acc@0.50	Acc@0.25	Acc@0.50
1	49.74	46.25	78.31	71.74
2	51.08	47.17	<u>78.21</u>	72.15
3	51.28	47.58	<u>78.21</u>	<u>72.35</u>
4	51.39	<u>47.89</u>	77.90	72.56
5	52.00	48.10	77.60	72.05
6	<u>51.70</u>	48.10	77.49	71.84

Table 1: Ablation study on the recurrent localization steps m of ReCOT in terms of Acc@0.25(%) $\uparrow$ and Acc@0.50(%) $\uparrow$ on the test set of CVOGL-DetGeo dataset. Bold and Underline indicate the best and second-best results, respectively.

Component	Ground→Satellite		Drone→Satellite	
	Acc@0.25	Acc@0.50	Acc@0.25	Acc@0.50
w/o RFEM	46.45	42.24	50.15	46.04
w/o $\mathcal{L}_{\text{SAM}}$	49.74	44.81	70.91	65.78
w/o $\mathcal{L}_{\text{Token}}$	50.57	47.58	72.05	66.80
ReCOT (Ours)	52.00	48.10	78.21	72.35

Table 2: Ablation study on the components of ReCOT in terms of Acc@0.25(%) $\uparrow$ and Acc@0.50(%) $\uparrow$ on the test set of CVOGL-DetGeo dataset.

RFEM	Ground→Satellite		Drone→Satellite	
	Acc@0.25	Acc@0.50	Acc@0.25	Acc@0.50
w/o $\mathbf{M}$	51.18	47.48	75.44	70.30
Replace $\mathbf{F}_{lr}$ with $\mathbf{F}_{hr}$	48.00	44.19	75.85	70.09
Replace $\mathbf{F}_{hr}$ with $\mathbf{F}_{lr}$	47.49	43.28	75.64	67.11
ReCOT (Ours)	52.00	48.10	78.21	72.35

Table 3: Ablation study inside the RFEM in terms of Acc@0.25(%) $\uparrow$ and Acc@0.50(%) $\uparrow$ on the test set of CVOGL-DetGeo dataset.

Fig. 2). The auxiliary loss $\mathcal{L}_{\text{aux}}$ is the sum of $\mathcal{L}_{\text{Token}_i}$ across all recurrent steps and the SAM-based distillation loss $\mathcal{L}_{\text{SAM}}$. The balancing coefficient α is set to 1 in all experiments.

4 Experiment

Datasets

CVOGL-DetGeo dataset (Sun et al. 2023) divides the task into “Ground → Satellite” task and “Drone → Satellite” task. It contains 6,239 pairs of “Ground → Satellite” view and 6,239 pairs of “Drone → Satellite” view query and reference images. The ground view query images, drone view query images, and satellite view reference images are sized to 512×256 , 256×256 , and 1024×1024 , respectively. Each cross-view task uses 4,343 pairs for training, 923 pairs for validation, and 973 pairs for testing. Our experiments utilize the training set to train the model and select the best model on the validation set for evaluation on the test set.

Method	Ground→Satellite				Drone→Satellite			
	Test Set		Validation Set		Test Set		Validation Set	
	Acc@0.25	Acc@0.50	Acc@0.25	Acc@0.50	Acc@0.25	Acc@0.50	Acc@0.25	Acc@0.50
CVM-Net (Hu et al. 2018)	4.73	0.51	5.09	0.87	20.14	3.29	20.04	3.47
RK-Net (Lin et al. 2022)	7.40	0.82	8.67	0.98	19.22	2.67	19.94	3.03
L2LTR (Yang, Lu, and Zhu 2021)	10.69	2.16	12.24	1.84	38.95	6.27	38.68	5.96
Polar-SAFA (Shi et al. 2019)	20.66	3.19	19.18	2.71	37.41	6.58	36.19	6.39
TransGeo (Zhu, Shah, and Chen 2022)	21.17	2.88	21.67	3.25	35.05	6.47	34.78	5.42
SAFA (Shi et al. 2019)	22.20	3.08	20.59	3.25	37.41	6.58	36.19	6.39
GeoDTR+ (Zhang et al. 2024b)	14.19	5.14	14.08	1.95	16.03	4.73	15.71	3.68
DetGeo (Sun et al. 2023)	45.43	42.24	46.70	43.99	61.97	57.66	59.81	55.15
VaGeo (Li et al. 2025)	48.21	45.22	47.56	44.42	66.19	61.87	64.25	59.59
OCGNet (Huang et al. 2025)	<u>51.49</u>	<u>47.69</u>	48.54	<u>44.20</u>	<u>68.39</u>	<u>63.93</u>	<u>66.52</u>	<u>61.86</u>
ReCOT (Ours)	52.00	48.10	<u>48.43</u>	43.66	78.21	72.35	74.00	67.17

Table 4: Comparisons in terms of Acc@0.25(%) $\uparrow$ and Acc@0.50(%) $\uparrow$ on the CVOGL-DetGeo dataset. Bold and Underline indicate the best and second-best results, respectively.

Method	Param (M) $\downarrow$	FPS $\uparrow$
DetGeo (Sun et al. 2023)	<u>73.8</u>	29.5
VaGeo (Li et al. 2025)	—	—
OCGNet (Huang et al. 2025)	74.8	<u>27.7</u>
ReCOT (Ours)	29.9	25.7

Table 5: Model complexity and runtime comparison in terms of parameters (M) $\downarrow$ and FPS $\uparrow$. Bold and Underline indicate the best and second-best results, respectively.

Evaluation Metrics

Following the pioneering work DetGeo (Sun et al. 2023), we adopt Acc@0.25 and Acc@0.50 for evaluation. For each pair of the query and reference image, we select the box with the highest confidence output by our model as the final prediction box. Additionally, we use the parameter to show the model efficiency. Higher Acc@0.25, higher Acc@0.50, and fewer parameter denote better performance. Please refer to the supplementary material for detailed introduction of evaluation metrics.

Implementation Details

We conduct all experiments on four NVIDIA GeForce RTX 4090 GPUs, with implementations based on PyTorch (Paszke et al. 2019). For training, we adopt the AdamW (Loshchilov and Hutter 2017) as the optimizer and set the initial learning rate to 0.0025, the weight decay rate to 0.0001, and the batch size to 16. We train our network for 300 epochs for all the experiments. Since CVOGL is a relatively new task, we follow the pioneering work (Sun et al. 2023) and select five CVIGL approaches (Hu et al. 2018; Lin et al. 2022; Yang, Lu, and Zhu 2021; Shi et al. 2019; Zhu, Shah, and Chen 2022) and the existing CVOGL approaches (Sun et al. 2023; Li et al. 2025; Huang et al. 2025) as our comparison methods. The results of CVIGL comparison methods can be found in previous works (Sun et al. 2023; Huang et al. 2025; Li et al. 2025). Additionally, we

adopt swin transformer (Swin-t) (Liu et al. 2021) as the image encoder for its superior performance on various fields (Zhang et al. 2024a; Shi et al. 2024; Chi, Yuan, and Wang 2023; Zamir et al. 2022). We set the hyper-parameter m to 6 during training.

Ablation Study

Effect of Total Recurrent Localization Steps. Table 1 investigates the impact of varying the total number of recurrent steps m on ReCOT performance. For the Ground→Satellite task, increasing m from 1 (one-shot prediction) to 5 yields consistent improvements, with the best performance achieved at $m = 5$. In contrast, for the Drone→Satellite task, the performance saturates earlier, with $m = 3$ giving the relatively optimal results. This discrepancy indicates that the optimal number of recurrent steps is task-dependent due to differences in viewpoint variations and feature alignment difficulty across scenarios. Further increasing m beyond the optimal point does not bring additional gains and may even slightly degrade performance, which can be attributed to over-refinement and overfitting in deeper iterations (Hur and Roth 2019; Cao et al. 2023; Yu et al. 2023). Based on the results, we set m to 5 for the Ground→Satellite task, while $m = 3$ for the Drone→Satellite task. We keep this setting in other experiments of this work.

Effect of Components in ReCOT. Table 2 presents the ablation study of the key components in ReCOT on the CVOGL-DetGeo test set. Removing the RFEM module (w/o RFEM) leads to a noticeable performance drop, confirming that early reference feature enhancement is critical for guiding tokens to focus on relevant regions. Similarly, removing the SAM-based distillation loss $\mathcal{L}_{\text{SAM}}$ (w/o $\mathcal{L}_{\text{SAM}}$) results in performance degradation, indicating the importance of accurate prompt semantics understanding. The full ReCOT model achieves the best performance across both scenarios. In addition, removing the token-guidance loss $\mathcal{L}_{\text{Token}}$ (w/o $\mathcal{L}_{\text{Token}}$) also degrades performance, which highlights its role in encouraging the learnable tokens to accurately capture object-relevant areas during recurrent refinement. These



Figure 5: Visualization of some representative results produced by our ReCOT and the previous work (Sun et al. 2023). Please refer to the zoomed-in view for better visualization.

results collectively validate that RFEM, $\mathcal{L}_{\text{SAM}}$, and $\mathcal{L}_{\text{Token}}$ complement each other, contributing to the robust performance of ReCOT.

Effect of Components Within the RFEM. Table 3 investigates the contributions of components within the RFEM. Removing the weight matrix $\mathbf{M}$ (w/o $\mathbf{M}$) leads to a noticeable performance drop, particularly in the Drone→Satellite setting. This confirms that the spatial attention stage benefit the reference feature enhancement. Furthermore, replacing the low-resolution semantic feature $\mathbf{F}_{\text{lr}}$ with the high-resolution feature $\mathbf{F}_{\text{hr}}$ and replacing $\mathbf{F}_{\text{hr}}$ with $\mathbf{F}_{\text{lr}}$ both result in performance degradation, highlighting the importance of leveraging multi-resolution reference features in RFEM. Please refer to the supplementary material for more ablation study results.

Comparison Results

Quantitative Results. Table 4 compares ReCOT with existing SOTA cross-view localization approaches on the CVOGL-DetGeo dataset. ReCOT consistently outperforms all competitors in both the Ground→Satellite and Drone→Satellite settings, achieving new SOTA performance across almost all metrics in the test set. Despite relatively lower performance on the validation set for Ground→Satellite, it maintains the best performance on the test set, indicating stronger generalization ability. Notably, as shown in Table 5, these performance gains are achieved with significantly fewer parameters, representing a 60% reduction in model size. ReCOT is a little slower in inference speed compared to the previous CVOGL works (Sun et al. 2023; Huang et al. 2025) due to iterative framework (Cao et al. 2023). However, it still achieves a competitive inference speed, making ReCOT both efficient and scalable for real-world applications. Compared to the Drone→Satellite setting, the improvement of ReCOT on Ground→Satellite is relatively smaller. This is mainly due to the larger viewpoint

gap and severe occlusions in ground-view images (Sun et al. 2023), where objects are often partially visible or obstructed by surrounding structures. Moreover, ground images typically contain more background clutter, making recurrent localization harder. We believe integrating geometric priors (e.g., camera pose estimation) or multi-view fusion could further boost Ground→Satellite performance. Please refer to the supplementary material for more comparison results.

Qualitative Results. Fig. 5 visualizes some representative results. As the recurrent steps proceed, our model progressively refines the bounding boxes, leading to higher-quality localization compared to the single-shot prediction of the previous approach (Sun et al. 2023). However, the optimal number of recurrent steps varies across different scenarios, and excessive iterations may cause over-refinement (Hur and Roth 2019), resulting in a slight performance drop. Therefore, based on the ablation results in Table 1, we set the number of recurrent steps to 5 for Ground→Satellite and 4 for Drone→Satellite, which provides the best trade-off between accuracy and stability.

5 Conclusion

In this paper, we proposed ReCOT, which reformulates the CVOGL task as recurrent localization problem. By introducing learnable tokens to encode task-specific semantics and recurrently refine predictions, ReCOT addresses the limitations of one-shot detection paradigms of previous CVOGL works. We further incorporate a SAM-based knowledge distillation scheme to improve prompt understanding without incurring additional inference costs, and a RFEM to produce object-aware reference features via a hierarchical attention strategy. Extensive experiments on the CVOGL-DetGeo benchmark demonstrated that ReCOT achieves state-of-the-art performance with significantly fewer parameters and competitive inference speed.

Recurrent Cross-View Object Geo-Localization Supplementary Material

6 More Experimental Results

Detailed Introduction of Evaluation Metrics.

Definition of Acc@0.25 and Acc@0.50. The Acc@0.25 and Acc@0.50 measure the prediction accuracy under a specific intersection over union (IoU) threshold t between the predicted box b_p and ground box b_g as

$$\text{Acc}@t = \frac{1}{n} \sum_{i=1}^n \psi(t), \quad (8)$$

where

$$\psi(t) = \begin{cases} 1, \text{IoU}(b_p, b_g) \geq t \\ 0, \text{otherwise} \end{cases}, \quad (9)$$

$$\text{IoU}(b_p, b_g) = \frac{|b_p \cap b_g|}{|b_p \cup b_g|}. \quad (10)$$

More Ablation Study Results

Effect of the parameter α . Table 6 shows how the performance of ReCOT varies with different values of the balancing coefficient α , which controls the weight between the localization loss $\mathcal{L}_{\text{local}}$ and the auxiliary loss $\mathcal{L}_{\text{aux}}$. We observe that $\alpha = 1$ yields the best performance on both Ground→Satellite and Drone→Satellite. A smaller value ($\alpha = 0.1$) weakens the supervision of token alignment and SAM distillation, slightly reducing accuracy. Conversely, a larger value ($\alpha = 10$) overemphasizes auxiliary objectives, causing a notable drop. These results indicate that $\alpha = 1$ achieves the optimal trade-off.

Effect of the hyper-parameter n . Table 7 investigates the impact of the token number n on the performance of ReCOT. We observe a clear performance gain when increasing n from 1 to 100. This demonstrates that a sufficient number of tokens provides a richer representation of task-specific intent and better coverage of cross-view semantics, which benefits the recurrent refinement process. However, when n is further increased to 200, performance drops across all metrics. This decline is likely due to over-parameterization and token redundancy, which can introduce noise and hinder effective attention learning (Dosovitskiy et al. 2021), as also observed in token-scaling studies (Dosovitskiy et al. 2021; ?). Therefore, we set n to 100 in this work.

α	Ground→Satellite		Drone→Satellite	
	Acc@0.25	Acc@0.50	Acc@0.25	Acc@0.50
0.1	51.28	48.10	75.13	69.68
1	52.00	48.10	78.21	72.35
10	48.51	45.02	76.05	70.20

Table 6: Ablation study on the hyper-parameter α in terms of Acc@0.25(%) $\uparrow$ and Acc@0.50(%) $\uparrow$ on the test set of CVOGL-DetGeo dataset.

n	Ground→Satellite		Drone→Satellite	
	Acc@0.25	Acc@0.50	Acc@0.25	Acc@0.50
1	47.89	43.37	74.31	68.86
100	52.00	48.10	78.21	72.35
200	50.46	47.28	72.25	67.01

Table 7: Ablation study on the hyper-parameter n in terms of Acc@0.25(%) $\uparrow$ and Acc@0.50(%) $\uparrow$ on the test set of CVOGL-DetGeo dataset.

Method	Ground→Satellite		Drone→Satellite	
	Acc@0.25	Acc@0.50	Acc@0.25	Acc@0.50
DetGeo	45.43	42.24	61.97	57.66
DetGeo*	46.01	41.44	66.60	56.11
VaGeo	48.21	45.22	66.19	61.87
OCGNet	<u>51.49</u>	<u>47.69</u>	<u>68.39</u>	<u>63.93</u>
ReCOT (Ours)	52.00	48.10	78.21	72.35

Table 8: Comparisons between DetGeo (Sun et al. 2023), DetGeo using Swin-t (Liu et al. 2021) (DetGeo*), and our methods in terms of Acc@0.25(%) $\uparrow$ and Acc@0.50(%) $\uparrow$ on the test set of CVOGL-DetGeo dataset. Bold and Underline indicate the best and second-best results, respectively.

More Comparison Results.

Quantitative Results. As shown in Table 8, we replace the backbone of DetGeo (Sun et al. 2023) with Swin-t (Liu et al. 2021) for a more comprehensive and fair comparison. While the upgraded DetGeo* shows slight improvements on the Acc@0.25, it still lags behind our ReCOT by a large margin across all evaluation metrics. Moreover, the backbone replacement does not lead to consistent gains, as DetGeo* fails to outperform other recent CVOGL methods (Li et al. 2025; Huang et al. 2025) in Table 8. This demonstrates that the key factor of CVOGL does not lie in the backbone. In contrast, our proposed ReCOT achieves superior performance through a more effective and task-aligned framework, highlighting the importance of architectural innovations for CVOGL.

Qualitative Results. Fig. 6 provides more visualization results of the cross attention weights between the token and reference features during recurrent process. As the recurrent steps progress, object-relevant tokens gradually focus and strengthen around the expected object, enabling the bounding box to converge toward the correct location. In contrast, object-irrelevant tokens weaken and gradually stabilize over iterations. This behavior highlights the ability of ReCOT to dynamically disentangle object-focused information from irrelevant features, which is a key factor driving the success of our recurrent localization strategy.

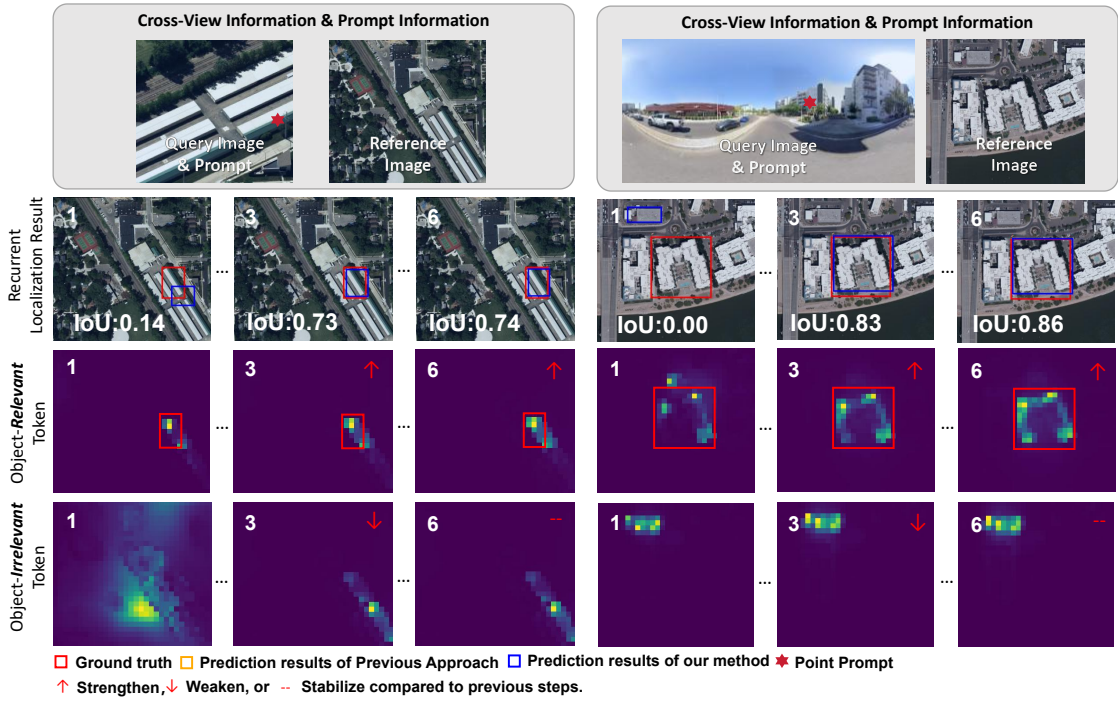


Figure 6: Visualizations of how our ReCOT works. The object-relevant tokens progressively focus and strengthen around the indicated object, while object-irrelevant tokens weaken and stabilize to background patterns.

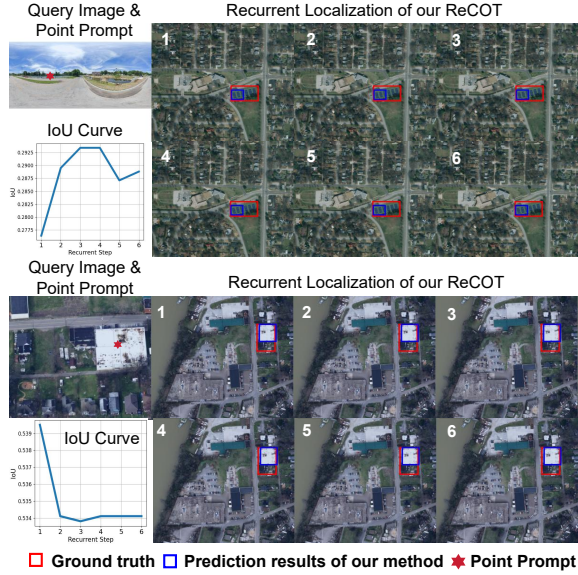


Figure 7: Examples of failure cases. Please refer to the zoomed-in view for better visualization.

7 Limitations and Future Work

As shown in Fig. 7, some failure cases of ReCOT are caused by ambiguous or imprecise point prompts, which often highlight only a part of the object rather than its entirety. This ambiguity misguides the token-driven recurrent refinement process, leading to suboptimal localization results. In the

future, we will incorporate multi-modal prompts (*e.g.*, textual descriptions or bounding boxes) to provide richer and more accurate prompt semantics. Such multi-modal guidance could reduce ambiguity and further improve the robustness and precision of recurrent localization.

References

- Cai, S.; Guo, Y.; Khan, S.; Hu, J.; and Wen, G. 2019. Ground-to-Aerial Image Geo-Localization With a Hard Exemplar Reweighting Triplet Loss. In *2019 IEEE/CVF International Conference on Computer Vision (ICCV)*, 8390–8399.
- Cao, S.-Y.; Hu, J.; Sheng, Z.; and Shen, H.-L. 2022. Iterative Deep Homography Estimation. In *2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 1869–1878.
- Cao, S.-Y.; Zhang, R.; Luo, L.; Yu, B.; Sheng, Z.; Li, J.; and Shen, H.-L. 2023. Recurrent Homography Estimation Using Homography-Guided Image Warping and Focus Transformer. In *2023 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 9833–9842.
- Carion, N.; Massa, F.; Synnaeve, G.; Usunier, N.; Kirillov, A.; and Zagoruyko, S. 2020. End-to-End Object Detection with Transformers. *arXiv e-prints arXiv:2005.12872*.
- Chi, K.; Yuan, Y.; and Wang, Q. 2023. Trinity-Net: Gradient-Guided Swin Transformer-Based Remote Sensing Image Dehazing and Beyond. *IEEE Transactions on Geoscience and Remote Sensing*, 61: 1–14.

- Chini, M.; Pierdicca, N.; and Emery, W. J. 2009. Exploiting SAR and VHR Optical Images to Quantify Damage Caused by the 2003 Bam Earthquake. *IEEE Transactions on Geoscience and Remote Sensing*, 47(1): 145–152.
- Deuser, F.; Habel, K.; and Oswald, N. 2023. Sample4Geo: Hard Negative Sampling For Cross-View Geo-Localisation. In *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, 16847–16856.
- Dosovitskiy, A.; Beyer, L.; Kolesnikov, A.; Weissenborn, D.; Zhai, X.; Unterthiner, T.; Dehghani, M.; Minderer, M.; Heigold, G.; Gelly, S.; Uszkoreit, J.; and Houlsby, N. 2021. An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale. *arXiv e-prints arXiv:2010.11929*.
- Edstedt, J.; Sun, Q.; Bökman, G.; Wadenbäck, M.; and Felsberg, M. 2024. RoMa: Robust Dense Feature Matching. *IEEE Conference on Computer Vision and Pattern Recognition*.
- Guo, Y.; Choi, M.; Li, K.; Boussaid, F.; and Bennamoun, M. 2022. Soft Exemplar Highlighting for Cross-View Image-Based Geo-Localization. *IEEE Transactions on Image Processing*, 31: 2094–2105.
- Hu, S.; Feng, M.; Nguyen, R. M. H.; and Lee, G. H. 2018. CVM-Net: Cross-View Matching Network for Image-Based Ground-to-Aerial Geo-Localization. In *2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 7258–7267.
- Huang, Z.; Aryal, J.; Nahavandi, S.; Lu, X.; Lim, C. P.; Wei, L.; and Zhou, H. 2025. Object-level Cross-view Geo-localization with Location Enhancement and Multi-Head Cross Attention.
- Hur, J.; and Roth, S. 2019. Iterative Residual Refinement for Joint Optical Flow and Occlusion Estimation.
- Kirillov, A.; Mintun, E.; Ravi, N.; Mao, H.; Rolland, C.; Gustafson, L.; Xiao, T.; Whitehead, S.; Berg, A. C.; Lo, W.-Y.; Dollar, P.; and Girshick, R. 2023. Segment Anything. In *2023 IEEE/CVF International Conference on Computer Vision (ICCV)*, 4015–4026.
- Kumar, A.; Kim, H.; and Hancke, G. P. 2013. Environmental Monitoring Systems: A Review. *IEEE Sensors Journal*, 13(4): 1329–1339.
- Lentsch, T.; Xia, Z.; Caesar, H.; and Kooij, J. F. P. 2023. SliceMatch: Geometry-Guided Aggregation for Cross-View Pose Estimation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 17225–17234.
- Li, Z.; Yuan, X.; Liu, W.; and Xu, X. 2025. VAGeo: View-specific Attention for Cross-View Object Geo-Localization. *ArXiv*, 2501.07194.
- Lin, J.; Zheng, Z.; Zhong, Z.; Luo, Z.; Li, S.; Yang, Y.; and Sebe, N. 2022. Joint Representation Learning and Keypoint Detection for Cross-View Geo-Localization. *IEEE Transactions on Image Processing*, 31: 3780–3792.
- Liu, L.; and Li, H. 2019. Lending Orientation to Neural Networks for Cross-View Geo-Localization. In *2019 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 5617–5626.
- Liu, Z.; Lin, Y.; Cao, Y.; Hu, H.; Wei, Y.; Zhang, Z.; Lin, S.; and Guo, B. 2021. Swin Transformer: Hierarchical Vision Transformer using Shifted Windows. In *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*.
- Loshchilov, I.; and Hutter, F. 2017. Fixing Weight Decay Regularization in Adam. *arXiv e-prints arXiv:1711.05101*.
- Lu, X.; Li, Z.; Cui, Z.; Oswald, M. R.; Pollefeys, M.; and Qin, R. 2020. Geometry-Aware Satellite-to-Ground Image Synthesis for Urban Areas. In *2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 856–864.
- Lu, X.; Luo, S.; and Zhu, Y. 2022. It’s Okay to Be Wrong: Cross-View Geo-Localization With Step-Adaptive Iterative Refinement. *IEEE Transactions on Geoscience and Remote Sensing*, 60: 1–13.
- Milletari, F.; Navab, N.; and Ahmadi, S.-A. 2016. V-Net: Fully Convolutional Neural Networks for Volumetric Medical Image Segmentation. *arXiv e-print arXiv:1606.04797s*.
- Paszke, A.; Gross, S.; Massa, F.; Lerer, A.; Bradbury, J.; Chanan, G.; Killeen, T.; Lin, Z.; Gimelshein, N.; Antiga, L.; Desmaison, A.; Kopf, A.; Yang, E.; DeVito, Z.; Raison, M.; Tejani, A.; Chilamkurthy, S.; Steiner, B.; Fang, L.; Bai, J.; and Chintala, S. 2019. PyTorch: An Imperative Style, High-Performance Deep Learning Library. In Wallach, H.; Larochelle, H.; Beygelzimer, A.; d’Alché-Buc, F.; Fox, E.; and Garnett, R., eds., *Advances in Neural Information Processing Systems*, volume 32.
- Regmi, K.; and Shah, M. 2019. Bridging the Domain Gap for Ground-to-Aerial Image Matching. In *2019 IEEE/CVF International Conference on Computer Vision (ICCV)*, 470–479.
- Sarlin, P.-E.; DeTone, D.; Yang, T.-Y.; Avetisyan, A.; Straub, J.; Malisiewicz, T.; Buló, S. R.; Newcombe, R.; Kotschieder, P.; and Balntas, V. 2023. OrienterNet: Visual Localization in 2D Public Maps with Neural Matching. In *2023 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 21632–21642.
- Shi, L.; Chen, Y.; Liu, M.; and Guo, F. 2024. DuST: Dual Swin Transformer for Multi-modal Video and Time-Series Modeling. In *2024 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) Workshops*, 4537–4546.
- Shi, Y.; and Li, H. 2022. Beyond Cross-view Image Retrieval: Highly Accurate Vehicle Localization Using Satellite Image. In *2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 16989–16999.
- Shi, Y.; Liu, L.; Yu, X.; and Li, H. 2019. Spatial-Aware Feature Aggregation for Image based Cross-View Geo-Localization. In *Advances in Neural Information Processing Systems*, volume 32.
- Shi, Y.; Yu, X.; Campbell, D.; and Li, H. 2020a. Where Am I Looking At? Joint Location and Orientation Estimation by Cross-View Matching. In *2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 4063–4071.

- Shi, Y.; Yu, X.; Liu, L.; Campbell, D.; Koniusz, P.; and Li, H. 2022. Accurate 3-DoF Camera Geo-Localization via Ground-to-Satellite Image Matching. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 45: 2682–2697.
- Shi, Y.; Yu, X.; Liu, L.; Zhang, T.; and Li, H. 2020b. Optimal Feature Transport for Cross-View Image Geo-Localization. *Proceedings of the AAAI Conference on Artificial Intelligence*, 34(07): 11990–11997.
- Singamaneni, P. T.; Bachiller-Burgos, P.; Manso, L. J.; Garrell, A.; Sanfeliu, A.; Spalanzani, A.; and Alami, R. 2024. A survey on socially aware robot navigation: Taxonomy and future challenges. *The International Journal of Robotics Research*, 43(10): 1533–1572.
- Sun, Y.; Ye, Y.; Kang, J.; Fernandez-Beltran, R.; Feng, S.; Li, X.; Luo, C.; Zhang, P.; and Plaza, A. 2023. Cross-View Object Geo-Localization in a Local Region With Satellite Imagery. *IEEE Transactions on Geoscience and Remote Sensing*, 61: 1–16.
- Teed, Z.; and Deng, J. 2020. RAFT: Recurrent All-Pairs Field Transforms for Optical Flow.
- Toker, A.; Zhou, Q.; Maximov, M.; and Leal-Taixe, L. 2021. Coming Down to Earth: Satellite-to-Street View Synthesis for Geo-Localization. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 6488–6497.
- Wang, X.; Xu, R.; Cui, Z.; Wan, Z.; and Zhang, Y. 2023. Fine-Grained Cross-View Geo-Localization Using a Correlation-Aware Homography Estimator. *ArXiv*, abs/2308.16906.
- Woo, S.; Park, J.; Lee, J.-Y.; and Kweon, I. S. 2018. CBAM: Convolutional Block Attention Module. In *Proc. Eur. Conf. Comput. Vis.*, 3–19. Cham: Springer Nature Switzerland.
- Yang, H.; Lu, X.; and Zhu, Y. 2021. Cross-view Geo-localization with Layer-to-Layer Transformer. In *Advances in Neural Information Processing Systems*, volume 34, 29009–29020. Curran Associates, Inc.
- Yao, J.; Hong, D.; Gao, L.; and Chanussot, J. 2022. Multi-modal Remote Sensing Benchmark Datasets for Land Cover Classification. In *2022 IEEE International Geoscience and Remote Sensing Symposium (IGARSS)*, 4807–4810.
- Yu, J.; Chang, J.; He, J.; Zhang, T.; Yu, J.; and Wu, F. 2023. Adaptive Spot-Guided Transformer for Consistent Local Feature Matching. In *2023 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 21898–21908.
- Zamir, S. W.; Arora, A.; Khan, S.; Hayat, M.; Khan, F. S.; and Yang, M.-H. 2022. Restormer: Efficient Transformer for High-Resolution Image Restoration. In *2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 5728–5739.
- Zhai, R.; Zou, J.; Gan, V. J.; Han, X.; Wang, Y.; and Zhao, Y. 2024. Semantic enrichment of BIM with IndoorGML for quadruped robot navigation and automated 3D scanning. *Automation in Construction*, 166: 105605.
- Zhang, T.; Zhang, X.; Zhu, X.; Wang, G.; Han, X.; Tang, X.; and Jiao, L. 2024a. Multistage Enhancement Network for Tiny Object Detection in Remote Sensing Images. *IEEE Transactions on Geoscience and Remote Sensing*, 62: 1–12.
- Zhang, X.; Li, X.; Sultani, W.; Chen, C.; and Wshah, S. 2024b. GeoDTR: Toward Generic Cross-View Geolocalization via Geometric Disentanglement. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 46(12): 10419–10433.
- Zhang, X.; Zhang, X.; Cao, S.-Y.; Yu, B.; Zhang, C.; and Shen, H.-L. 2024c. MRF³Net: An Infrared Small Target Detection Network Using Multireceptive Field Perception and Effective Feature Fusion. *IEEE Transactions on Geoscience and Remote Sensing*, 62: 1–14.
- Zhu, S.; Shah, M.; and Chen, C. 2022. TransGeo: Transformer Is All You Need for Cross-View Image Geo-Localization. In *2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 1162–1171.
- Zhu, Y.; Sun, B.; Lu, X.; and Jia, S. 2022. Geographic Semantic Network for Cross-View Image Geo-Localization. *IEEE Transactions on Geoscience and Remote Sensing*, 60: 1–15.