

Communication-Efficient Relative Pose Estimation with Vision Foundation Models for Ephemeral Collaborative Perception

Qihang Li^{1,†}, Jo-Hao Huang^{1,†}, Jiewen Liu¹, Suyoung Kang², Hao Zhang², and Peng Gao¹

Abstract—Relative pose estimation is a fundamental capability for collaborative perception and coordination in multi-robot systems. However, robots encountering each other in real-world environments often operate in short interaction windows and must operate under limited communication bandwidth with intermittent or missing visual overlap caused by occlusions or limited fields of view. Existing approaches typically rely on global reference frames, assume sustained view overlap, or incur prohibitive communication costs, thereby limiting their applicability to ephemeral collaborative perception. To address these challenges, we introduce *communication-efficient relative pose estimation* (CERPE), a system-level framework that coordinates vision foundation models to jointly estimate ego-motion and inter-robot relative pose. CERPE reduces unnecessary raw-observation exchange by using continuously shared fixed-size descriptors to gate event-triggered raw-image requests independently of pose estimation. Non-overlapping encounters are handled by propagating inter-robot relative poses through metrically scaled ego-motion, thus maintaining relative pose estimates even in the absence of visual overlap. Experiments in simulation and real-world robots show that CERPE improves 6-DoF relative pose estimation over selected baselines in ephemeral collaborative perception.

I. INTRODUCTION

Multi-robot systems have been extensively studied over the past decades due to their ability to enable parallel task execution, expand sensing coverage, and improve robustness through redundancy. They have been deployed across a wide range of real-world applications, including connected autonomous driving [1], multi-robot-assisted warehouse logistics [2], and distributed environmental monitoring [3]. To enable effective multi-robot systems, collaborative perception plays a critical role in establishing shared environmental awareness among robot teammates to support downstream tasks such as coordinating actions and improving perception robustness.

To enable collaborative perception, a fundamental capability is reliable relative pose estimation among robot teammates, which determines the 6-DoF pose of one robot with respect to another using only onboard sensors, without relying on global references such as GPS, motion capture systems, or prebuilt maps. As shown in Fig. 1, during an ephemeral encounter at an urban intersection, two connected vehicles must estimate their relative pose within a short interaction window before fusing observations for collaborative perception. Within this brief encounter, the connected vehicles operate under limited communication bandwidth and may experience weak or even no visual overlap due to occlusions and limited fields of

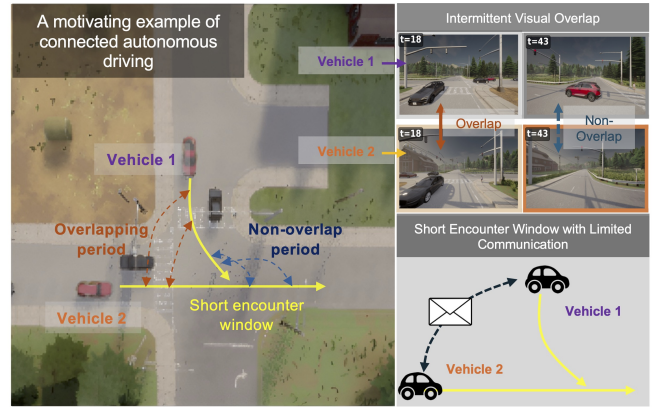


Fig. 1. Motivating scenario for communication-efficient relative pose estimation in connected autonomous driving. When two connected vehicles encounter each other at an urban intersection, they must estimate their relative pose within a short encounter window to merge observations for collaborative perception, without relying on global references. This requires addressing two key challenges: (1) intermittent visual overlap, which can lead to weak or absent overlapping observations, and (2) limited communication bandwidth, which constrains information exchange during the brief encounter.

view. These factors make relative pose estimation particularly challenging in real-world scenarios.

Given the importance of relative pose estimation, several approaches have been developed. Global reference-based methods rely on GPS, indoor tracking systems, or maps, but these references are often unreliable or unavailable [4]. Collaborative SLAM (CSLAM) methods recover full robot trajectories through optimization over loop closures [5]–[7], but they are computationally expensive and require maintaining a global map. Vision-based registration methods estimate relative pose through feature matching or geometric alignment [8]–[13], yet performance drops under large viewpoint changes. Recently, vision foundation models have demonstrated strong generalization in feed-forward 3D reconstruction, enabling camera pose estimation by implicitly inferring scene geometry and view transformations from visual observations [14]–[18]. In multi-robot systems, however, using such representations for inter-robot place recognition or pose estimation would require transmitting raw observations or high-dimensional visual features. This communication cost, together with the common assumption of overlapping observations, limits their effectiveness for short-duration or non-overlapping encounters.

To address these challenges, we propose communication-efficient relative pose estimation (CERPE), which organizes vision foundation models into a training-free pipeline for

[†]Equal contribution. ¹North Carolina State University, Raleigh, NC 27695, USA. ²University of Massachusetts Amherst, Amherst, MA 01002, USA. Emails: {qli39, pgao5}@ncsu.edu.

metrically scaled ego-motion and inter-robot relative pose estimation under communication constraints during ephemeral collaborative perception. Each robot uses SALAD [19] to encode observations into fixed-size descriptors and shares them with its local pose in lightweight messages. Descriptor similarity serves as an online proxy for sufficient visual overlap. For ego-motion estimation, each robot continuously runs VGGT [15] over the current observation and a bounded keyframe memory selected by temporal spacing and descriptor-based viewpoint change. Gate activation triggers raw-observation exchange and cross-robot VGGT inference for direct relative pose estimation. We resolve VGGT's monocular scale ambiguity by aligning its predicted depth with metric depth from Metric3Dv2 [20]. Otherwise, the latest relative pose is propagated using metrically scaled ego-motion from both robots.

The primary contribution of this paper is CERPE, a conditional execution pipeline that gates raw-observation exchange by descriptor similarity, uses metric depth to calibrate monocular scale, and switches to SE(3) ego-motion propagation when visual overlap disappears. Specifically,

- CERPE leverages vision foundation models to enable two-stage information sharing with fixed-size descriptor messages for metrically scaled ego-motion and inter-robot relative pose estimation, while explicitly addressing intermittent and missing visual overlap in multi-robot collaborative perception.
- CERPE supports online deployment for metric-scale 6-DoF relative pose estimation in small multi-robot teams, without relying on global references such as GPS or maps. In our evaluated settings, CERPE reduces position error by up to **72%** and rotation error by up to **90%** relative to selected baselines across both overlapping and non-overlapping observations.

II. RELATED WORK

A. Collaborative Perception

Collaborative perception has received increasing attention in recent years due to its ability to overcome the limitations of single-robot perception, particularly in scenarios involving occlusions, long-range sensing, and restricted fields of view [21]. Collaborative perception has been widely studied for tasks, such as object localization [22], detection [1], trajectory forecasting [23], and connected vehicle control [24] in accident-prone scenarios. By associating observations from multiple viewpoints, these methods build a richer and more robust representation of the environment, leading to improved perception accuracy and reliability. However, most existing works implicitly assume that robots share sufficiently overlapping observations and can rely on accurate global references (e.g., GPS) to align multi-robot data. In the real-world multi-robot ephemeral encounter scenarios, visual overlap is often intermittent or weak, and global references may be unreliable or even unavailable.

B. Indirect Relative Pose Estimation

Indirect relative pose estimation refers to methods that first construct explicit intermediate representations such as feature correspondences or geometric alignment, and then recover the relative pose via geometric consistency constraints.

CSLAM-based methods estimate relative poses by maintaining a shared global map and recovering relative transformations through geometric consistency and loop closure constraints [3], [25], [26]. Although effective in structured environments, these approaches rely on persistent map consistency and reliable loop closures, which are difficult to ensure in large-scale or short-duration encounters, thereby limiting their applicability in ephemeral collaborative perception scenarios. Vision-based indirect methods register observations using feature correspondences [9], [13], [27] or geometric alignment [28], [29]. Recent advances incorporate diffusion-based refinement to improve correspondence quality [30], [31]. Despite improved accuracy, these methods typically assume sufficient visual overlap and degrade under large viewpoint changes or intermittent observations.

C. Direct Relative Pose Estimation

Direct relative pose estimation refers to methods that directly predict relative poses from visual observations without constructing explicit intermediate representations. End-to-end regression approaches learn a direct mapping from image pairs to relative poses [12], [32]–[34]. However, they are constrained by their training distributions and typically require overlapping observations, limiting generalization to novel environments or non-overlapping scenarios.

More recently, feed-forward vision foundation models jointly predict scene geometry and camera poses in a single forward pass with strong generalization [14], [15], [18]. In particular, VGGT [15] accepts an arbitrary number of input frames, making it a suitable backbone for unifying visual odometry and collaborative perception within a single framework. However, these models predict geometry only up to an unknown scale, rely on high-dimensional features that are expensive to communicate, and assume overlapping observations. These limitations are particularly critical in ephemeral collaborative perception where interaction windows are short and communication bandwidth is constrained. Our method does not introduce a new foundation model; instead, it contributes a task-specific orchestration protocol over existing models, coupling descriptor-based overlap gating, selective memory insertion, direct inter-robot pose estimation, and SE(3) propagation for intermittent observations.

III. APPROACH

A. Problem Formulation

Consider a team of N robots operating in a shared environment. Without loss of generality, we represent a pair of robots as r_i and r_j , where $i, j \in \{1, \dots, N\}$ denote robot indices. Each robot r_i captures an RGB image sequence $\mathcal{I} = \{\mathbf{I}_i^1, \mathbf{I}_i^2, \dots\}$ along its trajectory, where $\mathbf{I}_i^t$ denotes the image captured by robot r_i at time t . Given a pair of image sequences $\mathcal{I}_i$ and $\mathcal{I}_j$ acquired by robot r_i and r_j , we aim

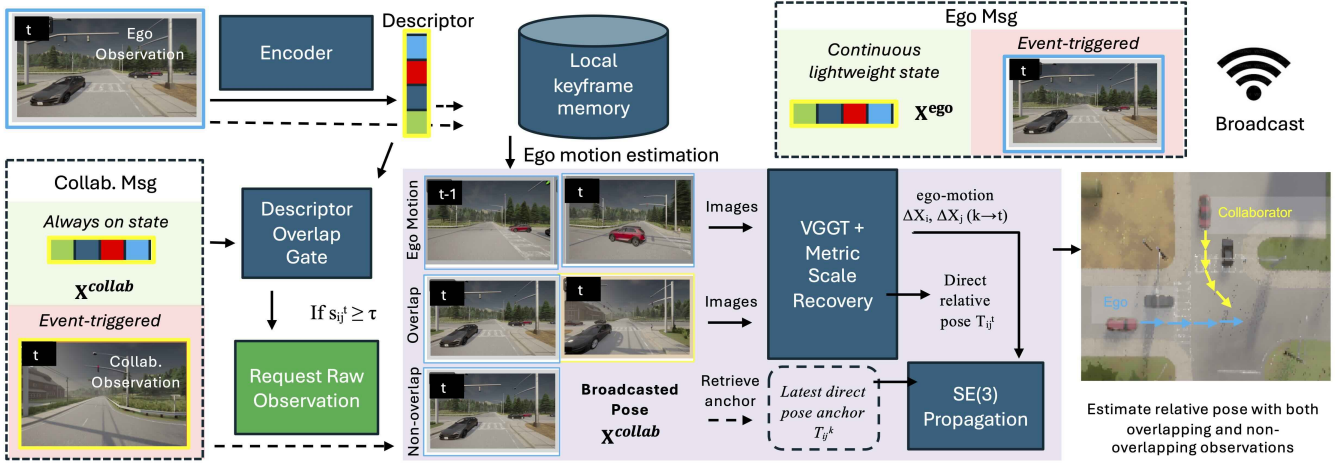


Fig. 2. Overview of our CERPE. Each robot continuously broadcasts its local pose and a fixed-size descriptor. Descriptor similarity gates raw-observation requests: gate activation triggers direct inter-robot pose estimation; otherwise, CERPE propagates the latest direct pose anchor using both robots' metrically scaled ego-motion.

to estimate the relative pose $T_{ij}^t \in \text{SE}(3)$ between robots r_i and r_j at every time step t . We decompose the relative pose estimation ϕ into two sub-problems, including

- *Ego-motion estimation.* Each robot maintains a bounded keyframe memory $\mathcal{M}_i^t$ and estimates its local pose X_i^t by processing the current observation together with this memory. Since all poses predicted in a VGGT forward pass share the same local coordinate frame, the robot recovers incremental ego-motion by composing the corresponding local poses, e.g., $T_i^{t,t-1} = (X_i^t)^{-1} X_i^{t-1}$.
- *Relative pose estimation.* When robot r_i and robot r_j have visual overlap at time t , the relative pose T_{ij}^t can be directly estimated from their overlapping observations (I_i^t, I_j^t) based on vision foundation models [15]. In non-overlapping periods, the relative pose is propagated using a previously established relative pose estimate and each robot's accumulated ego-motion, which is defined as $T_{ij}^t = (X_j^t)^{-1} X_j^{t-k} T_{ij}^{t-k} (X_i^{t-k})^{-1} X_i^t$, where k denotes the most recent time step when overlapping observations are available.

We address two key challenges of this formulation: (1) enabling vision foundation models with metric scale awareness for ego-motion and relative pose estimation, and (2) overcoming the communication constraints imposed by high-dimensional visual features, thereby enabling relative pose estimation under both overlapping and non-overlapping observations.

B. Metrically Scaled Ego-Motion Estimation

To estimate the ego-motion $T_i^{t,t-1} = (X_i^t)^{-1} X_i^{t-1}$, we adopt VGGT [15] as the pose estimation backbone, denoted f_{vggt} . At each time step t , robot r_i runs the backbone on the current image together with its local memory queue $\mathcal{M}_i^{t-1}$:

$$\{X_i^k, D_i^k\}_{I_i^k \in \mathcal{M}_i^{t-1} \cup \{I_i^t\}} = f_{\text{vggt}}(\mathcal{M}_i^{t-1} \cup \{I_i^t\}). \quad (1)$$

The predicted poses are expressed in a common local coordinate frame induced by the current memory. Therefore,

when the previous frame or keyframe is present in the memory, the ego-motion is recovered by composing the two local poses as $T_i^{t,t-1} = (X_i^t)^{-1} X_i^{t-1}$; otherwise the robot composes motions through the most recent stored keyframe. However, as the vision foundation model infers geometry from monocular observations without access to metric depth, the predicted rotations are scale-invariant, whereas the translation component of X_i^t is defined only up to an unknown scale due to the inherent scale ambiguity of monocular vision. The memory is updated using both temporal spacing and descriptor-based viewpoint change. Let c_i^t denote the cosine similarity between the current descriptor and the descriptor of the most recent local keyframe, and let Δ_i^t denote their index gap. The current observation is selected as a new keyframe when

$$(c_i^t < \tau_m \wedge \Delta_i^t > \Delta_k) \vee c_i^t < \tau_k, \quad (2)$$

where τ_m selects frames with moderate viewpoint change after sufficient temporal spacing, and the stricter threshold $\tau_k < \tau_m$ immediately accepts frames with large viewpoint change. The selected keyframes are saved into a bounded memory queue $\mathcal{M}_i^t = \{I_k\}_{k=1}^K$, where K denotes the total number of keyframes in memory.

Given the memory queue $\mathcal{M}_i^t$, we recover the scale factor $s \in \mathbb{R}^+$ for the entire memory by comparing them against metric depth references D_k^* , which are obtained from a metric depth foundation model Metric3Dv2 [20]. Formally, we compute the metric depth reference as

$$D_k^* = f_{\text{m3d}}(I_k), \quad I_k \in \mathcal{M}_i^t \quad (3)$$

where f_{m3d} denotes the metric depth foundation model. Then we compute the scale factor as

$$s(\mathcal{M}_i^t) = \frac{1}{K} \sum_{k=1}^K \text{median}_p \frac{D_k^*(p)}{D_k(p)}. \quad (4)$$

where $D_k(p)$ denotes the depth value at pixel p in the k -th keyframe. $\frac{D_k^*(p)}{D_k(p)}$ provides a local estimate of the scale at

each pixel. By taking the median over all pixels across all keyframes in the memory, we obtain a stable scale estimate s . The scale factor is then applied to recover the metrically scaled ego-motion as

$$\mathbf{T}_i^{t,t-1} = \mathbf{S}_i \cdot (\mathbf{X}_i^t)^{-1} \mathbf{X}_i^{t-1}, \quad \mathbf{S}_i = \text{diag}(s_i, s_i, s_i, 1) \quad (5)$$

where $\mathbf{S}_i \in \text{Sim}(3)$ scales the pose from the up-to-scale output to metrically consistent ego-motion.

C. Communication-Efficient Relative Pose Estimation

To estimate the relative pose $\mathbf{T}_{ij}^t$ between robot r_i and robot r_j , we use the same VGGT model as backbone and memory queue to store keyframes. Once the robot r_i receives the observations sent from its collaborator robot r_j , the relative pose between them is computed as

$$\mathbf{T}_{ij}, \mathbf{D}_i, \mathbf{D}_j = f_{\text{vggt}}(\mathbf{I}_i, \mathbf{I}_j). \quad (6)$$

where $\mathbf{T}_{ij}$ denotes the relative pose between robot r_i and r_j , and $\mathbf{D}_i, \mathbf{D}_j$ denote their respective depth estimates.

However, $\mathbf{I}_i$ and $\mathbf{I}_j$ may not overlap due to limited fields of view or occlusion, causing direct estimation to fail. We therefore use descriptor similarity as an online proxy for sufficient visual overlap. Formally, given an image $\mathbf{I}_i$, we utilize the same encoder $f_{\text{vggt}}^{\text{enc}}$ used in VGGT to extract a dense feature map $\mathbf{F}_i = f_{\text{vggt}}^{\text{enc}}(\mathbf{I}_i)$. This dense feature map encodes rich semantic and geometric information but is high-dimensional, with spatial resolution proportional to the input image size and channel dimensionality on the order of hundreds. As a result, directly sharing $\mathbf{F}_i$ between robots would incur prohibitive communication cost and is impractical for bandwidth-constrained multi-robot collaboration.

We therefore adopt two-stage communication: robots frequently exchange fixed-size descriptors to gate candidate overlap, while gate activation triggers raw-observation sharing for inter-robot pose estimation. Specifically, we adopt the visual aggregation head proposed in SALAD [19], which compresses the dense visual feature map into a compact global descriptor by optimally aggregating local patch features through Sinkhorn-based optimal transport [35], while discarding non-informative regions via a learned dustbin mechanism. The compressed descriptor is computed as

$$\mathbf{d}_i = \psi(\mathbf{F}_i) \in \mathbb{R}^D, \quad (7)$$

where ψ denotes the visual aggregation head. For fixed D and precision, message size is fixed, while the compression ratio depends on image resolution and encoding. The robots compute cosine similarity $\text{score} = \frac{\mathbf{d}_i^\top \mathbf{d}_j}{\|\mathbf{d}_i\| \|\mathbf{d}_j\|}$. A pair with $\text{score} \geq \tau$ is treated as a candidate overlap and triggers raw-observation exchange; otherwise the gate remains closed.

If the descriptor gate identifies a candidate overlap, the ego robot requests the raw observation from its collaborator. Once the collaborator's observation is received, it is inserted into the ego robot's memory queue, and the inter-robot relative pose is estimated using Eq. (6). If the gate is not activated, the ego robot does not request the raw observation, thereby avoiding unnecessary communication and computation. The

Algorithm 1 Robot r_i egocentric relative pose estimation.

Require: Observation $\mathbf{I}_i^t$; messages $\{\mathbf{m}_j^t\}$; memory $\mathcal{M}_i^{t-1}$
Ensure: Relative poses $\{\mathbf{T}_{ij}^t\}$; ego-motion $\mathbf{X}_i^t$
Feature Extraction & Communication:
1: $\mathbf{d}_i^t \leftarrow \psi(f_{\text{vggt}}^{\text{enc}}(\mathbf{I}_i^t))$ $\triangleright$ Fixed-size descriptor
2: Broadcast $\mathbf{m}_i^t = (i, \mathbf{X}_i^t, \mathbf{d}_i^t)$
3: **for** each received $\mathbf{m}_j^t = (j, \mathbf{X}_j^t, \mathbf{d}_j^t)$ **do**
4: $s_{ij}^t \leftarrow \frac{\mathbf{d}_i^{t\top} \mathbf{d}_j^t}{\|\mathbf{d}_i^t\| \|\mathbf{d}_j^t\|}$ $\triangleright$ Descriptor similarity
5: **if** $s_{ij}^t \geq \tau$ **then**
6: Request $\mathbf{I}_j^t$; enqueue into $\mathcal{M}_i^t$
7: **end if**
8: **end for**
Pose Estimation & Scale Recovery:
9: $\{\mathbf{X}_k, \mathbf{D}_k\} \leftarrow f_{\text{vggt}}(\mathcal{M}_i^t \cup \{\mathbf{I}_i^t\})$
10: $s \leftarrow \frac{1}{K} \sum_k \text{median}_p \mathbf{D}_k^*(p) / \mathbf{D}_k(p)$
11: $\mathbf{X}_k \leftarrow s \cdot \mathbf{X}_k, \forall k$ $\triangleright$ Metrically scaled poses
Relative Pose Estimation:
12: **for** each robot r_j **do**
13: **if** $s_{ij}^t \geq \tau$ **then** $\triangleright$ Candidate overlap
14: $\mathbf{T}_{ij}^t \leftarrow (\mathbf{X}_j^t)^{-1} \mathbf{X}_i^t$
15: **else** $\triangleright$ Non-overlapping
16: $\mathbf{T}_{ij}^t \leftarrow (\mathbf{X}_j^t)^{-1} \mathbf{X}_j^{t-k} \mathbf{T}_{ij}^{t-k} (\mathbf{X}_i^{t-k})^{-1} \mathbf{X}_i^t$
17: **end if**
18: **end for**
Memory Management:
19: $c_i^t \leftarrow s(\mathbf{d}_i^t, \mathbf{d}_i^{t_{\text{last}}})$; $\Delta_i^t \leftarrow t - t_{\text{last}}$
20: **if** $(c_i^t < \tau_m \wedge \Delta_i^t > \Delta_k) \vee c_i^t < \tau_k$ **then**
21: Enqueue $\mathbf{I}_i^t$ into $\mathcal{M}_i^t$ $\triangleright$ New keyframe
22: **end if**
23: **if** $|\mathcal{M}_i^t| > L$ **then**
24: Evict oldest entry from $\mathcal{M}_i^t$
25: **end if**

relative pose is propagated using a previously established relative pose estimate $\mathbf{T}_{ij}^{t-k}$ at time $t - k$ and two robots' accumulated ego-motion, which is computed as

$$\mathbf{T}_{ij}^t = (\mathbf{X}_j^t)^{-1} \mathbf{X}_j^{t-k} \mathbf{T}_{ij}^{t-k} (\mathbf{X}_i^{t-k})^{-1} \mathbf{X}_i^t \quad (8)$$

where k denotes the most recent time step when overlapping observations are available. To estimate relative pose in both overlapping and non-overlapping conditions, directly sharing the raw image $\mathbf{I}$ among robots is bandwidth-intensive. Since overlap occurs only intermittently, CERPE triggers this high-bandwidth transmission only when the descriptor-based gate is activated, reducing unnecessary raw-observation exchange.

Algorithm 1 summarizes the decentralized execution. Each robot broadcasts its fixed-size descriptor and local pose, requests raw observations only upon gate activation, and otherwise propagates the latest inter-robot estimate using metric ego-motion. Bounded memory retains frames after sufficient temporal and viewpoint change, or immediately after a large viewpoint change.

IV. EXPERIMENTS

A. Experiment Setup

We comprehensively evaluate our CERPE in three experimental settings:

- A high-fidelity connected autonomous driving (CAD) simulation based on CARLA [36], which generates connected vehicle trials across 4 intersections, comprising 71 encounter scenarios with a total of 3,461 image pairs.
- Real-world CAD trials built upon OpenMars [37], which captures multi-robot long-range encounters in urban environments. We use the multi-robot subset and select 23 trajectory pairs (6,900 image pairs) that contain camera-level overlap, as many sequences in this dataset have weak visual overlap between robots' views.
- A physical robot team of three Agilx Limo Pro robots deployed in 13 manually constructed indoor scenarios that cover large perspective changes, non-overlapping views, occluded views, and dynamic objects. We uniformly select 130 keyframes from 13 trajectories, yielding 382 pairwise image pairs with valid ground-truth poses. Ground-truth poses are obtained from a motion capture system.

In all experiments, we adopt VGGT [15] as the pose estimation backbone with training-free metric scale recovery defined in Eq. (4) and pretrained compact descriptors from SALAD [19] for place recognition defined in Eq. (7). In simulation, metric scale is recovered via Metric3Dv2 [20]. In real-world deployments, onboard depth sensors provide metric depth for scale alignment. The keyframe memory uses $\tau_m = 0.9$ for moderate viewpoint change, $\tau_k = 0.2$ for large viewpoint change, a fixed minimum index gap Δ_k , the descriptor-gate threshold $\tau = 0.2$, and the memory capacity $L = 5$.

For evaluation, we compare CERPE with three previous methods and one baseline, including

- **SuperGlue** [9] that is a modular feature-matching pipeline that detects and matches sparse keypoints across image pairs. We recover the pose via RANSAC [38]. As SuperGlue cannot recover metric scale, we report its translation performance using ground-truth scale.
- **NOPE** [39]: an object-level graph matching method that estimates relative pose by exploiting spatial relationships among detected objects (e.g., vehicles and pedestrians) in the scene.
- **Co-VisNet** [12]: an end-to-end model that directly regresses robot relative pose from image pairs.
- A baseline **CERPE w/o EM** that performs relative pose estimation without using ego-motion propagation to address non-overlapping observations.

To quantitatively evaluate and compare CERPE with other methods, we use three metrics, including position error, rotation error, and success rate. Given the estimated translation $\hat{\mathbf{t}}$ and ground-truth translation $\mathbf{t}$, and the estimated and ground-truth rotations represented as unit quaternions $\hat{\mathbf{q}}, \mathbf{q} \in \mathbb{S}^3$, the

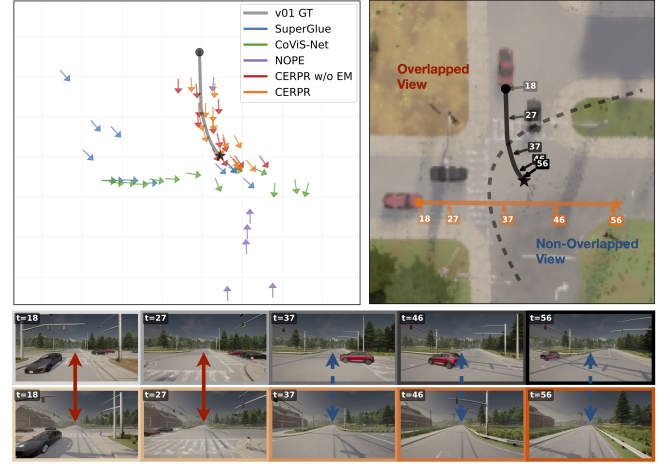


Fig. 3. Qualitative results in CAD simulation. **Right:** bird's-eye view of both vehicles' ground truth trajectories, with the dashed line separating overlapping and non-overlapping periods. **Left:** ground-truth and predicted relative poses of the target vehicle (black trajectory) in the reference vehicle's frame (orange trajectory). **Bottom:** temporally corresponding image pairs from the two vehicles with red arrows and blue dashed arrows indicating overlapping and non-overlapping views, respectively.

position and rotation errors are defined as:

$$e_{\text{pos}} = \|\hat{\mathbf{t}} - \mathbf{t}\|_2, \quad (9)$$

$$e_{\text{rot}} = 4 \arcsin\left(\frac{1}{2} \min(\|\hat{\mathbf{q}} - \mathbf{q}\|, \|\hat{\mathbf{q}} + \mathbf{q}\|)\right) \quad (10)$$

where the rotation metric follows the geodesic distance on $\mathbb{S}^3$ via Roma [40]. We additionally report the success rate, defined as the percentage of frames for which a method produces a valid pose estimate. Failures arise when a method cannot recover a pose (e.g., insufficient feature matches or inliers) or when the method determines that the input is outside its operating range (e.g., no detectable overlap).

B. Evaluation in CAD Simulation

We first evaluate our approach in CAD simulation where two connected vehicles encounter each other at an urban intersection. The encounter window lasts approximately 3–5 seconds and involves large relative viewpoint changes. As the vehicles traverse the intersection, occlusions caused by surrounding traffic frequently lead to transitions from overlapping to non-overlapping visual observations.

As shown in Fig. 3, the ego vehicle initially shares overlapping views with its collaborator, but the overlap gradually disappears as the vehicles diverge. NOPE predicts relative poses in a nearly opposite direction to the ground truth. This is because two visually similar black vehicles appear in both robots' views. Since NOPE relies on object-level spatial relationships to estimate relative poses, it confuses object identities under such ambiguous visual cues. Co-VisNet shows a strong bias toward a fixed formation prior, consistently predicting its collaborator at a short distance in front of the ego vehicle, regardless of the true relative pose under large viewpoint changes and partial occlusions. In contrast, our method maintains temporally stable relative pose estimates

TABLE I

QUANTITATIVE RESULTS OF OUR CERPE COMPARED WITH SELECTED BASELINES BASED ON POSITION, ROTATION ERROR AND SUCCESS RATE.

Method	CAD Simulation			Real-world CAD			Physical Robot Teams		
	e_{pos} (m) ↓	e_{rot} (rad) ↓	Succ. (%) ↑	e_{pos} (m) ↓	e_{rot} (rad) ↓	Succ. (%) ↑	e_{pos} (m) ↓	e_{rot} (rad) ↓	Succ. (%) ↑
SuperGlue [9]	16.00 [†]	0.895	77.5	7.62 [†]	0.169	93.7	0.12 [†]	0.165	98.4
NOPE [39]	9.41	0.776	33.4	39.41	1.682	45.0	—	—	—
Co-VisNet [12]	16.22	0.906	100	29.95	0.271	100	0.88	0.257	100
CERPE w/o EM	5.03/3.51[†]	0.128	100	10.27/5.84[†]	0.063	100	0.45/0.09[†]	0.096	100
CERPE	<u>8.66/5.06[†]</u>	0.089	100	8.26/2.94[†]	0.063	100	<u>0.54/0.11[†]</u>	0.086	100

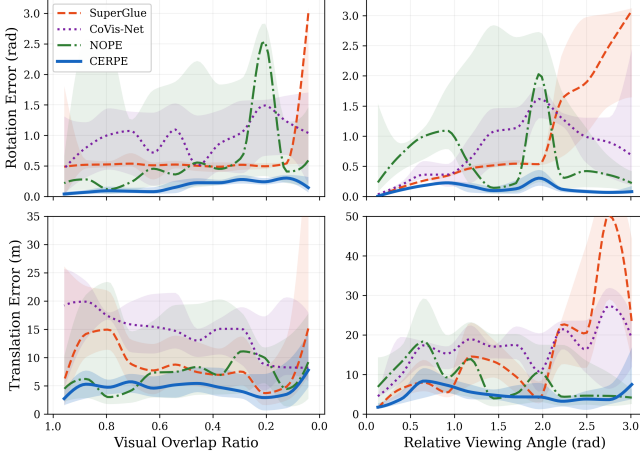
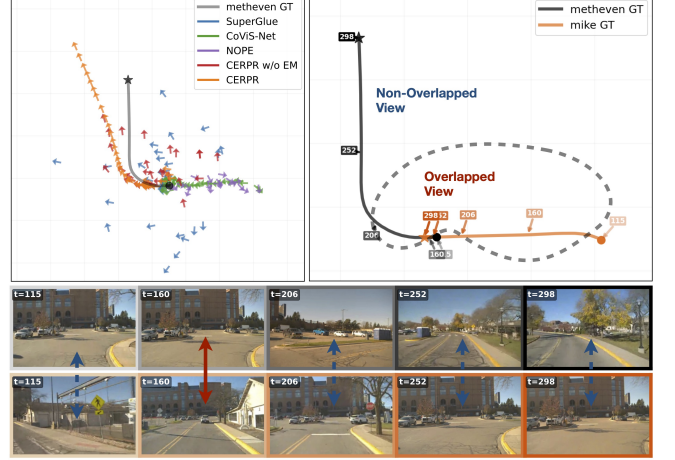
[†] Translation with ground-truth scale. For CERPE, both predicted and GT-scale translations are shown (pred./GT).

Fig. 4. Analysis of performance under varying degrees of visual overlap and viewpoint changes. Compared with selected baselines, our approach remains stable across large variations in visual overlap and view angles.

throughout the encounter by jointly reasoning over multiple keyframes stored in memory, enabling continued estimation after visual overlap is lost.

The quantitative results are shown in Table I. Our method achieves competitive results across the evaluated scenarios while consistently maintaining a 100% success rate. Among the comparison methods, NOPE achieves a translation error of 9.41 m but succeeds on only 33.4% of pairs. Co-VisNet always produces an estimate but with substantially higher errors; and SuperGlue fails on a notable fraction of pairs due to insufficient feature matches. Our baseline CERPE w/o EM outperforms the comparison methods in CAD simulation, achieving 5.03 m position error and 0.128 rad rotation error. Our full approach further improves rotation to 0.089 rad but increases translation error to 8.66 m. This result highlights an operating condition of ego-motion propagation: in high-overlap settings, direct estimates already provide strong translation constraints, while short propagation windows can add ego-motion and scale drift. Thus, propagation is best viewed as a continuity mechanism for missing-overlap periods rather than a module that uniformly improves every metric.

Fig. 4 analyzes how each method’s performance varies with visual overlap ratio and relative viewing angle given the CAD simulation. Our method maintains a median rotation error below 0.15 rad across all overlap bins, whereas SuperGlue’s median exceeds 1.5 rad when overlap falls below 20%,

Fig. 5. Qualitative results in the real-world CAD. **Right:** bird’s-eye view of both vehicles’ trajectories with the gray dashed line separating overlapping and non-overlapping periods. **Left:** ground-truth and predicted relative poses. **Bottom:** temporally corresponding image pairs.

reflecting its dependence on sufficient feature correspondences. NOPE exhibits an interquartile range spanning over 1 rad in rotation even at moderate overlap, indicating unstable object-level matching. Co-VisNet achieves comparable rotation accuracy to ours at high overlap but incurs translation errors exceeding 15 m due to its lack of metric scale. Regarding viewing angle, the comparison methods degrade sharply beyond 2 rad relative rotation, while our method remains more stable even at viewing angles approaching π rad, indicating improved stability in geometrically challenging regimes.

C. Evaluation in Real-World CAD

We further evaluate CERPE in the real-world CAD scenario. Unlike in simulation, where visual overlap between vehicles can be deliberately controlled, real-world driving scenarios typically exhibit much smaller overlap due to unconstrained vehicle motion, long inter-vehicle distances, and occlusions from surrounding traffic.

As shown in Fig. 5, when a low-overlap encounter occurs, our method recovers a lower-error relative pose than the comparison methods by incorporating the inter-robot image into its memory. Moreover, after the overlap ends, our approach maintains temporal continuity through ego-motion propagation combined with the previously established esti-

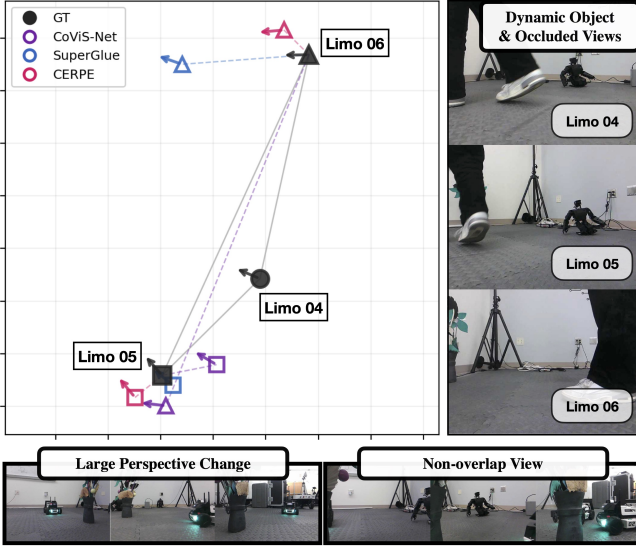


Fig. 6. Qualitative results on physical robot teams in indoor scenarios with large perspective changes, non-overlapping views, and occlusions.

mate, whereas the other methods either fail entirely or degrade to unconstrained guesses. Notably, although error propagation during non-overlapping periods inevitably introduces some drift, the incorporation of visual odometry improves the temporal consistency of the estimated trajectory, preventing abrupt jumps exhibited by baselines that treat each frame independently. When overlap resumes, the direct estimate corrects accumulated drift.

The quantitative results of real-world CAD are shown in Table I. Compared with simulation, this scenario features longer trajectories with extended non-overlapping periods. Our CERPE achieves lower rotation error of 0.063 rad than all comparison methods, with the nearest competitor SuperGlue at 0.169 rad with ground-truth scale. In translation, our CERPE achieves the best predicted-scale result of 8.26 m, followed by our baseline CERPE w/o EM at 10.27 m. With ground-truth scale, ours further reduces translation error to 2.94 m, compared with SuperGlue at 7.62 m with a 93.7% success rate.

D. Evaluation on Physical Robot Teams

To evaluate physical-robot pose estimation beyond simulation and driving scenarios, we deploy CERPE on a team of three Agilex Limo Pro robots in indoor environments. Unlike the CAD scenarios where robots move along road-constrained trajectories, the indoor robot scenes feature close-range interactions with frequent transitions between overlapping and non-overlapping views, occluded observations caused by surrounding objects, and dynamic disturbances. As shown in Fig. 6, our CERPE predictions (red) are closer to the reference poses (black) for all three robots, whereas Co-VisNet (purple) and SuperGlue (blue) exhibit notable deviations, particularly for the distant Limo 06 where large perspective changes and dynamic occlusions are present.

As shown in Table I, NOPE fails entirely on this dataset, as the indoor scenes lack the structured object categories (e.g.,

vehicles, pedestrians) that its object-level matching relies on. SuperGlue achieves the lowest translation error of 0.12 m with ground-truth scale and succeeds on 376 out of 382 pairs at 98.4%, benefiting from the close-range setting where sufficient feature correspondences are available. However, it still fails on 6 pairs and cannot recover metric scale. Co-VisNet produces estimates for all pairs but with substantially higher errors of 0.88 m position and 0.257 rad rotation, indicating limited generalization to the indoor domain. Our CERPE achieves the lowest rotation error of 0.086 rad, outperforming SuperGlue at 0.165 rad by 48% and Co-VisNet at 0.257 rad by 67%. In translation, CERPE w/o EM achieves the best predicted-scale result of 0.45 m and ground-truth-scale result of 0.09 m, while maintaining 100% success. These results suggest that our approach transfers from driving scenarios to close-range indoor robot settings.

Finally, we run an online deployment test on 2–3 robots to check execution, communication, and controller integration; because motion-capture coverage does not span the full operating region, quantitative physical-robot accuracy is evaluated on the sampled mocap-covered scenarios above rather than from this closed-loop test. All computation runs on a laptop with an Intel i9-14900HX CPU and an Nvidia RTX 4080 GPU, and robots communicate over a shared wireless network using 640×480 JPEG-compressed images. The descriptor-only message is 16.9 KB; overlap-triggered image exchange adds a compressed image averaging 26 KB in our setup. The full CERPE runs in 438 ms for two robots and 780 ms for three robots, while a 50 Hz low-level controller consumes the latest pose estimate between perception updates.

V. CONCLUSION

We introduced CERPE, a communication-efficient framework for metric relative pose estimation under ephemeral encounters. CERPE coordinates existing vision foundation models through descriptor-gated image exchange, metric scale recovery, and ego-motion propagation from the latest direct pose anchor. Across CARLA, OpenMars, and indoor robot evaluations, CERPE maintains a 100% pose-estimation success rate and improves accuracy over selected baselines in several settings without relying on GPS or maps. The results also show that propagation preserves continuity during missing-overlap intervals but can accumulate scale and translation drift when direct overlap constraints are strong. Direct inter-robot estimates therefore act as intermittent geometric anchors, while bounded-memory ego-motion bridges the intervals between them; descriptor screening remains separate from pose inference, so raw RGB transmission occurs only after gate activation.

Several limitations remain. Scale quality directly affects propagation and relative pose consistency, motivating additional metric cues such as stereo, LiDAR, or inertial sensing. The current shared-compute prototype targets small teams; larger teams require lighter or distributed inference. Finally, CERPE has no back-end optimization, and factor-graph refinement over keyframes and encounter anchors could reduce long-term drift.

REFERENCES

- [1] R. Xu, H. Xiang, Z. Tu, X. Xia, M.-H. Yang, and J. Ma, “V2x-vit: Vehicle-to-everything cooperative perception with vision transformer,” in *Proceedings of the European Conference on Computer Vision (ECCV)*, 2022.
- [2] J. Li, A. Tinka, S. Kiesel, J. W. Durham, T. K. S. Kumar, and S. Koenig, “Lifelong multi-agent path finding in large-scale warehouses,” *Proceedings of the AAAI Conference on Artificial Intelligence*.
- [3] P. Gao, J. Liang, Y. Shen, S. Son, and M. C. Lin, “Visual, spatial, geometric-preserved place recognition for cross-view and cross-modal collaborative perception,” in *2023 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*. IEEE, 2023, pp. 11 079–11 086.
- [4] N. Gyagenda, J. V. Hatilima, H. Roth, and V. Zhmud, “A review of gnss-independent uav navigation techniques,” *Robotics and Autonomous Systems*, vol. 152, p. 104069, 2022.
- [5] T. Cieslewski, S. Choudhary, and D. Scaramuzza, “Data-efficient decentralized visual slam,” in *2018 IEEE international conference on robotics and automation (ICRA)*. IEEE, 2018, pp. 2466–2473.
- [6] Y. Tian, Y. Chang, F. H. Arias, C. Nieto-Granda, J. P. How, and L. Carlone, “Kimera-multi: Robust, distributed, dense metric-semantic slam for multi-robot systems,” *IEEE Transactions on Robotics*, vol. 38, no. 4, 2022.
- [7] H. Cao, S. Shreedharan, and N. Atanasov, “Multi-robot object slam using distributed variational inference,” *IEEE Robotics and Automation Letters*, 2024.
- [8] D. DeTone, T. Malisiewicz, and A. Rabinovich, “Superpoint: Self-supervised interest point detection and description,” in *Proceedings of the IEEE conference on computer vision and pattern recognition workshops*, 2018, pp. 224–236.
- [9] P.-E. Sarlin, D. DeTone, T. Malisiewicz, and A. Rabinovich, “Superglue: Learning feature matching with graph neural networks,” in *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2020, pp. 4938–4947.
- [10] P. Gao, R. Guo, H. Lu, and H. Z. Zhang, “Regularized graph matching for correspondence identification under uncertainty in collaborative perception,” in *Robotics science and systems*, 2021.
- [11] Z. Zhan, D. Gao, Y.-J. Lin, Y. Xia, and C. Wang, “ismatching: Imperative correspondence learning,” in *European Conference on Computer Vision*. Springer, 2024, pp. 183–200.
- [12] J. Blumenkamp, S. Morad, J. Gielis, and A. Prorok, “Covis-net: A cooperative visual spatial foundation model for multi-robot applications,” 2024.
- [13] S. Zhang and J. Ma, “Diffglue: Diffusion-aided image feature matching,” in *Proceedings of the 32nd ACM International Conference on Multimedia*, 2024, pp. 8451–8460.
- [14] S. Wang, V. Leroy, Y. Cabon, B. Chidlovskii, and J. Revaud, “Dust3r: Geometric 3d vision made easy,” in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2024, pp. 20 697–20 709.
- [15] J. Wang, M. Chen, N. Karaev, A. Vedaldi, C. Rupprecht, and D. Novotny, “Vggt: Visual geometry grounded transformer,” in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2025.
- [16] Y. Wang, J. Zhou, H. Zhu, W. Chang, Y. Zhou, Z. Li, J. Chen, J. Pang, C. Shen, and T. He, “ π^3 S: Permutation-equivariant visual geometry learning,” in *The Fourteenth International Conference on Learning Representations*, 2026.
- [17] N. Keetha, N. Müller, J. Schönberger, L. Porzi, Y. Zhang, T. Fischer, A. Knapitsch, D. Zauss, E. Weber, N. Antunes, J. Luiten, M. Lopez-Antequera, S. R. Bulò, C. Richardt, D. Ramanan, S. Scherer, and P. Kotschieder, “Mapanything: Universal feed-forward metric 3d reconstruction,” 2025. [Online]. Available: <https://arxiv.org/abs/2509.13414>
- [18] V. Leroy, Y. Cabon, and J. Revaud, “Grounding image matching in 3d with mast3r,” in *European conference on computer vision*. Springer, 2024, pp. 71–91.
- [19] S. Izquierdo and J. Civera, “Optimal transport aggregation for visual place recognition,” in *CVPR*, 2024.
- [20] M. Hu, W. Yin, C. Zhang, Z. Cai, X. Long, H. Chen, K. Wang, G. Yu, C. Shen, and S. Shen, “Metric3d v2: A versatile monocular geometric foundation model for zero-shot metric depth and surface normal estimation,” *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2024.
- [21] X. Gao, X. Zhang, Y. Lu, Y. Huang, L. Yang, Y. Xiong, and P. Liu, “A survey of collaborative perception in intelligent vehicles at intersections,” *IEEE Transactions on Intelligent Vehicles*, 2024.
- [22] Y.-C. Liu, J. Tian, C.-Y. Ma, N. Glaser, C.-W. Kuo, and Z. Kira, “Who2com: Collaborative perception via learnable handshake communication,” in *2020 IEEE International Conference on Robotics and Automation (ICRA)*. IEEE, 2020, pp. 6876–6883.
- [23] T.-H. Wang, S. Manivasagam, M. Liang, B. Yang, W. Zeng, and R. Urtasun, “V2vnet: Vehicle-to-vehicle communication for joint perception and prediction,” in *European conference on computer vision*. Springer, 2020, pp. 605–621.
- [24] J. Cui, H. Qiu, D. Chen, P. Stone, and Y. Zhu, “Coopernaut: End-to-end driving with cooperative perception for networked vehicles,” in *IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2022.
- [25] I. D. Miller, A. Cowley, R. Konkimalla, S. S. Shivakumar, T. Nguyen, T. Smith, C. J. Taylor, and V. Kumar, “Any way you look at it: Semantic crossview localization and mapping with lidar,” *IEEE Robotics and Automation Letters*, vol. 6, no. 2, pp. 2397–2404, 2021.
- [26] P. Gao and H. Zhang, “Long-term loop closure detection through visual-spatial information preserving multi-order graph matching,” in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 34, no. 06, 2020, pp. 10 369–10 376.
- [27] J. Sun, Z. Shen, Y. Wang, H. Bao, and X. Zhou, “Loftr: Detector-free local feature matching with transformers,” in *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2021, pp. 8922–8931.
- [28] H. Jiang, M. Salzmann, Z. Dang, J. Xie, and J. Yang, “Se (3) diffusion model-based point cloud registration for robust 6d object pose estimation,” *Advances in Neural Information Processing Systems*, vol. 36, pp. 21 285–21 297, 2023.
- [29] Z. Qin, H. Yu, C. Wang, Y. Guo, Y. Peng, and K. Xu, “Geometric transformer for fast and robust point cloud registration,” in *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2022, pp. 11 143–11 152.
- [30] J. Wang, C. Rupprecht, and D. Novotny, “Posediffusion: Solving pose estimation via diffusion-aided bundle adjustment,” in *Proceedings of the IEEE/CVF International Conference on Computer Vision*, 2023, pp. 9773–9783.
- [31] J. Nam, G. Lee, S. Kim, H. Kim, H. Cho, S. Kim, and S. Kim, “Diffusion model for dense matching,” in *The Twelfth International Conference on Learning Representations*, 2024. [Online]. Available: <https://openreview.net/forum?id=Zsfiqpf6K>
- [32] A. Kendall, M. Grimes, and R. Cipolla, “Posenet: A convolutional network for real-time 6-dof camera relocalization,” in *Proceedings of the IEEE international conference on computer vision*, 2015, pp. 2938–2946.
- [33] J. Y. Zhang, D. Ramanan, and S. Tulsiani, “Relpose: Predicting probabilistic relative rotation for single objects in the wild,” in *European Conference on Computer Vision*. Springer, 2022, pp. 592–611.
- [34] A. Lin, J. Y. Zhang, D. Ramanan, and S. Tulsiani, “Relpose++: Recovering 6d poses from sparse-view observations,” in *2024 International Conference on 3D Vision (3DV)*. IEEE, 2024, pp. 106–115.
- [35] M. Cuturi, “Sinkhorn distances: Lightspeed computation of optimal transport,” *Advances in neural information processing systems*, vol. 26, 2013.
- [36] A. Dosovitskiy, G. Ros, F. Codevilla, A. Lopez, and V. Koltun, “CARLA: An open urban driving simulator,” in *Proceedings of the 1st Annual Conference on Robot Learning*, 2017, pp. 1–16.
- [37] Y. Li, Z. Li, N. Chen, M. Gong, Z. Lyu, Z. Wang, P. Jiang, and C. Feng, “Multiagent multitaversal multimodal self-driving: Open mars dataset,” in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2024, pp. 22 041–22 051.
- [38] M. A. Fischler and R. C. Bolles, “Random sample consensus: a paradigm for model fitting with applications to image analysis and automated cartography,” *Communications of the ACM*, vol. 24, no. 6, pp. 381–395, 1981.
- [39] H. Huang, D. Xu, H. Zhang, and P. Gao, “Non-overlap-aware egocentric pose estimation for collaborative perception in connected autonomy,” 2025.
- [40] R. Bréquier, “Deep regression on manifolds: a 3D rotation case study,” 2021.