

Think While You Map: Asynchronous Vision-Language Agents for Incremental 3D Scene Graphs

Deniz Bickici^{1,3}, Michael Pabst^{1,3}, Shohei Mori¹, Dieter Schmalstieg^{1,2}

¹ University of Stuttgart, Germany

² Graz University of Technology, Austria

³ IMPRS-IS, Germany

`deniz.bickici@visus.uni-stuttgart.de`

Abstract. Open-vocabulary 3D scene graph methods typically operate in two stages: first reconstruct, then enrich with vision-language models, leaving the graph unqueryable during exploration. We argue that this sequential coupling is unnecessary and propose an asynchronous architecture in which lightweight online mapping runs concurrently with heavyweight semantic refinement. A probabilistic voxel-based backbone maintains stable object identities incrementally, while background VLM agents progressively enrich the graph. This framework resolves duplicate object tracks through semantic loop closure, attaches fine-grained visual attributes and derives spatial relations between objects. A multi-target frame scheduler amortizes VLM cost by selecting a small set of informative frames that jointly cover multiple targets. The resulting scene graph is queryable during exploration and grows in semantic richness over time. Our method matches or outperforms existing open-vocabulary 3D scene graph methods on semantic segmentation (ScanNet, Replica) and surpasses the prior state-of-the-art across three visual grounding benchmarks (Sr3D+, Nr3D, ScanRefer) by 15.3 to 18.8 A@0.25. Project page: <https://denizbickici.github.io/thinkgraphs/>

Keywords: 3D Scene Graphs · Open-Vocabulary · 3D Scene Understanding · Incremental Mapping · Visual Grounding

1 Introduction

Understanding the semantics of 3D environments is fundamental for embodied AI, robotics, and augmented reality. An agent must recognize objects, maintain their identities over time, infer relationships and support language-based queries for grounding and planning, e.g., “pick up the mug next to the laptop”. Scene graphs are a natural representation for this goal. They compactly encode entities as nodes and their interactions as edges, enabling downstream reasoning about “what is where” and “what is interacting with what.”

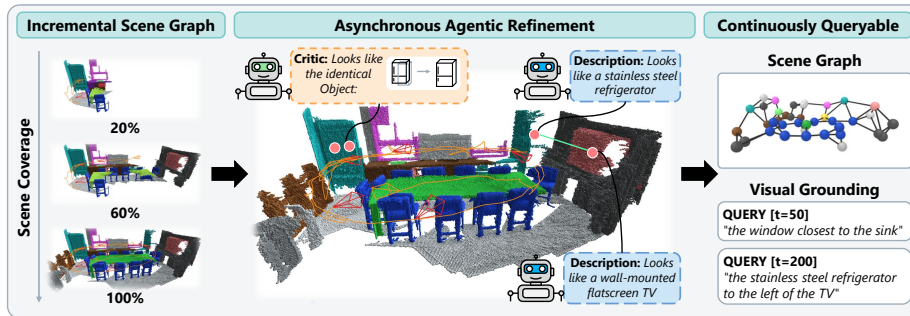


Fig. 1: ThinkGraphs decouples lightweight online mapping from heavyweight VLM reasoning, keeping the scene graph queryable during exploration. Asynchronous background agents refine the graph without blocking the mapping loop: a Critic Agent detects and merges fragmented object tracks, and a Description Agent injects fine-grained visual attributes, enabling complex grounding queries (e.g., “the stainless steel refrigerator left of the TV”) throughout exploration.

Recent work on 3D scene graphs has largely evolved along two complementary but currently disconnected directions: Semantic-rich offline reconstruction prevents querying unless the whole process is complete, while real-time reconstruction supports querying but only with limited semantic information. The former utilizes offline open-vocabulary systems and prioritizes rich labels, attributes, free-form descriptions, and functional relations [11, 19, 26, 45, 51], typically leveraging CLIP [32] embeddings and vision-language reasoning. These methods require completing instance-level reconstruction before enriching the scene with captions, labels, and inter-object relations. The latter are real-time mapping systems that incrementally update object graphs but rely on closed sets of semantics or shallow feature matching, trading the semantic richness of recent large language models (LLM) or vision-language models (VLM) for responsiveness [14, 24, 46].

The core technical challenge of combining these two paradigms lies in the instability of incremental tracking. Since objects are often fragmented or merged incorrectly under partial views and occlusions, it is impractical to query a VLM during exploration. Not only would this degrade the system’s responsiveness, it would also waste massive computational resources on analyzing duplicate or unstable object tracks. Consequently, previous systems have to wait for a fully consolidated offline reconstruction before applying VLM reasoning.

Many semantically rich pipelines already build their geometric map incrementally [11, 19]. What they lack is this combination of a robust enough 3D scene representation to make on-the-fly enrichment worthwhile, and an efficient scheduling mechanism to make it affordable.

Our key observation is that semantic reasoning does not need to wait for a stable graph; it can actively contribute to stabilizing it. This behavior is enabled by two factors. First, rather than the pairwise greedy association used by

prior methods, we accumulate per-voxel votes for each object track, producing substantially more robust object tracks, reducing the instability that motivated postponing semantics in the first place. Second, rather than brute-forcing VLM calls on every object at every frame, we treat VLM reasoning as a selective, asynchronous background process, one that not only enriches the graph with fine-grained attributes, but also actively repairs it by detecting and merging duplicate object tracks through semantic loop closure.

We therefore propose an asynchronous architecture in which a lightweight online mapper runs concurrently with two background VLM agents: a Critic Agent that detects and merges duplicate object tracks through semantic loop closure, and a Description Agent that attaches fine-grained visual attributes via a multi-target frame scheduler that amortizes VLM cost across informative frames. The resulting scene graph grows in semantic richness as exploration continues, yet never blocks online operation. Our method thus occupies the middle ground between offline open-vocabulary reconstruction and real-time closed-set mapping. It builds the graph incrementally during exploration, enriching semantics asynchronously rather than at frame rate.

To demonstrate the effectiveness of our approach, we extensively evaluate it on open-vocabulary semantic segmentation and three visual grounding benchmarks. In summary, we make the following contributions:

- We introduce *ThinkGraphs*, to our knowledge the first *asynchronous* architecture for open-vocabulary 3D scene graph construction. It decouples lightweight online mapping from heavyweight VLM reasoning, meaning the graph remains queryable during exploration and is continually enriched without blocking it.
- We propose a VLM-in-the-loop mechanism for semantic loop closure that detects and merges duplicate object tracks caused by incremental association drift, analogous to geometric loop closure in SLAM.
- We introduce a multi-target frame scheduler that selects a small number of informative frames for the Description Agent, each depicting multiple undescribed objects, so that a single VLM call enriches several nodes at once.
- We demonstrate that our incremental method outperforms all batch-based open-vocabulary scene graph methods on semantic segmentation and surpasses the prior state-of-the-art across three visual grounding benchmarks (Sr3D+, Nr3D, ScanRefer) by 15.3 to 18.8 A@0.25, showing that deferring all semantics to a post-hoc stage is unnecessary.

2 Related Work

Incremental 3D scene graphs. The incremental construction of 3D scene graphs focuses on updating structured spatial representations as new observations arrive. Kimera [37, 38] first demonstrated real-time metric-semantic mapping with hierarchical spatial layers, a design later extended by Hydra [14],

which pioneered incremental scene graphs by integrating signed distance fields with hierarchical graphs of objects, places, and rooms. S-Graphs+ [3] introduced a factor-graph formulation for hierarchical optimization, while SceneGraphFusion [47] fuses geometric over-segmentation with a Graph Neural Network (GNN) for incremental semantic prediction. Wu et al. [46] propose a sparse RGB-only variant that combines SLAM-based point tracking with multi-view feature aggregation. Further incremental methods either reuse temporal history [10] or incorporate knowledge from past observations [36] to enhance scene understanding. All of these methods remain restricted to closed-set vocabularies, limiting their applicability in open-world settings. More recently, Clio [24] introduced language interfaces into incremental mapping, averaging open-set features across temporal object tracks to form clusters relevant to the task, although it was limited to a predefined task list.

Compared to prior incremental scene graph mappers, our goal is to keep the graph open-vocabulary and queryable while exploration is ongoing. We achieve this goal by decoupling lightweight online tracking from heavyweight semantic enrichment: the mapper maintains stable object identities and geometry, while background VLM agents progressively refine labels and attributes without blocking the main loop; spatial relations are updated deterministically from 3D geometry.

Batch-processing of scene graphs. Many open-vocabulary 3D systems represent semantics as language-aligned features stored in dense maps or spatial memories [4, 12, 20, 30, 41, 48]. In contrast, batch scene graph pipelines commit these signals to discrete object nodes and relations, enabling compositional and relational queries.

ConceptGraphs [11] pioneered open-vocabulary 3D scene graphs by combining CLIP [32] embeddings with SAM2 [33] masks to construct 3D objects. Subsequent works extended this paradigm by adding hierarchical multi-level reasoning across rooms and floors [45], object affordances [51], and relational question answering [19]. Open3DSG [17] takes a learning-based approach by distilling 2D vision-language features into a graph neural network, producing language-aligned node and edge embeddings for open-vocabulary object and pairwise relationship prediction. PoVo [26] takes a different approach, operating on geometric super-points rather than per-frame masks and merging them into 3D instances through combined semantic and mask coherence.

While some of these methods build their geometric map incrementally [11, 19], all defer semantic enrichment to a post-hoc stage and therefore require a complete reconstruction before the graph becomes queryable. Our method instead exposes an intermediate graph throughout exploration and improves it incrementally: Semantic refinement is asynchronous and revisitable, so the system can answer queries early with coarse semantics and later update the same nodes or edges as better evidence emerges.

VLM reasoning in 3D scene understanding. Large vision-language models bring contextual reasoning capabilities that go beyond the embedding similarity used for object association and labeling. In open-vocabulary scene graphs, Con-

ceptGraphs [11] queries a VLM to caption each object node and an LLM to infer spatial edges, while BBQ [19] formulates object retrieval as deductive reasoning over a scene graph structure using an LLM.

For 3D visual grounding, LLM-Grounder [49] decompose complex natural-language queries with an LLM into sub-goals and iteratively invokes a 3D visual grounder (e.g., OpenScene [30] or LERF [16]) to localize targets in reconstructed scenes. SQA3D [23] benchmarks situated question answering that requires joint spatial and common-sense reasoning on 3D scans. Effective visual grounding in these pipelines often relies on visual prompting strategies such as Set-of-Mark [50], which overlays numeric tags on segmented image regions so that a VLM can refer to specific objects without fine-tuning.

Existing 3D mapping methods generally employ VLM/LLM reasoning in one of two ways: post-hoc enrichment after reconstruction, or synchronous query-time inference that blocks execution until the model responds. Our work targets the missing middle ground: We treat VLM reasoning as a background process that runs continuously, is scheduled selectively, and integrates results when available. This makes semantics available during mapping while controlling the computation via coverage-based frame selection rather than per-object per-frame querying.

Object-level association and merging. A key challenge in incremental scene graph mapping is maintaining consistent object identities as observations accumulate over time. Prior work addresses this through object-level SLAM [25], multi-object tracking and fusion [39, 40], and semantic object-based loop closure [18, 22]. However, association remains challenging in the presence of perceptual ambiguity and viewpoint changes.

In open-set mapping, ROMAN [31] aligns object submaps via graph-theoretic data association but assumes object nodes are already formed, while OpenVox [7] formulates incremental instance association as probabilistic voxel inference and Octree-Graph [44] proposes chronological group-wise merging with informative feature selection.

Most scene graph systems, whether incremental [14, 47] or batch-based [11, 26], rely on greedy pairwise association in learned embeddings (e.g., CLIP [32] or DINO [29]) and 3D geometric overlap. Such greedy strategies involve making early, local commitments based on limited evidence, which result in the accumulation of false negatives that lead to fragmented object tracks [14, 19]. BBQ [19] addresses this by consolidating duplicates in a periodic merge step, but repeated instances and viewpoint changes can still cause missed associations.

In contrast, our method couples probabilistic instance tracking with a VLM-based semantic loop closure mechanism that explicitly detects and corrects fragmented object tracks during online operation (Sec. 3.3).

3 Method

We introduce ThinkGraphs, an incremental open-vocabulary 3D scene graph pipeline that fuses geometric and visual evidence from sequentially posed RGB-

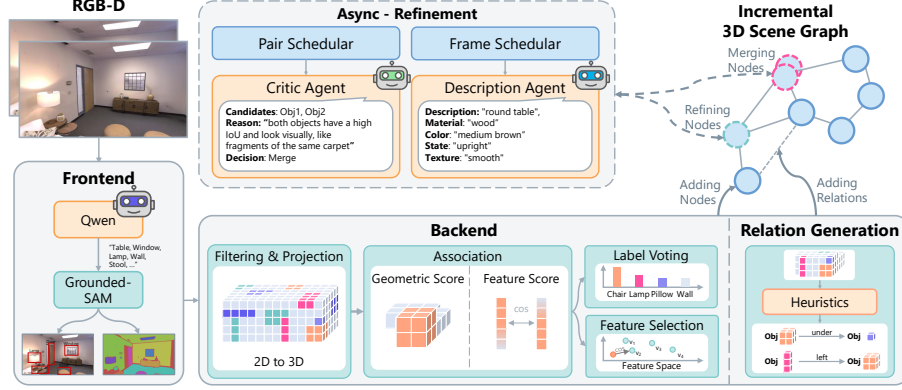


Fig. 2: Method Overview. (i) The frontend extracts grounded instances from each RGB-D frame (Sec. 3.1). (ii) The backend associates them into persistent 3D tracks with probabilistic voxel scoring and derives spatial edges deterministically from 3D geometry (Sec. 3.2). (iii) Two asynchronous VLM agents, a Critic Agent for semantic loop closure and a Description Agent for attribute enrichment, progressively refine the graph without blocking online operation (Sec. 3.3).

D frames $\{(I_t^{\text{rgb}}, I_t^{\text{depth}}, \theta_t)\}_{t=1}^N$ (Fig. 2), where θ_t is the camera pose. Our system incrementally constructs a scene graph $G_t = (\mathcal{V}_t, \mathcal{E}_t)$, where each node $v_i \in \mathcal{V}_t$ corresponds to a confirmed object track T_i , represented by a 3D point cloud $\mathbf{P}_i$, a selected CLIP visual embedding $\mathbf{f}_i^*$, a consensus label ℓ_i^* , and VLM-enriched attributes $\mathcal{A}_i$. Edges $\mathcal{E}_t$ encode spatial predicates derived from 3D geometry. Rather than requiring a complete reconstruction, the pipeline maintains a continuously queryable graph throughout exploration via three components: (i) a **frontend** that extracts grounded instance observations $\mathcal{O}_t$, (ii) a **backend** that lifts and associates them into persistent 3D object tracks, and (iii) **asynchronous VLM agents** that refine the scene graph without blocking the mapping loop.

3.1 Frontend for grounded instance extraction

Given an RGB image I_t^{rgb} at time t , we query Qwen3-VL-2B-Instruct [2] to produce a set of noun prompts $\mathcal{L}_t = \{\ell_j^t \mid j = 1, \dots, C_t\}$, as its context-aware prompts yield fewer missed detections than tag-based alternatives (Tab. 4). We feed these prompts into Grounded-SAM v2 [21, 33, 35], where Grounding-DINO [21] predicts text-conditioned boxes for each prompt, and SAM2 [33] then refines each box into a 2D instance mask. The frontend outputs a set of K_t grounded instances $\mathcal{O}_t = \{(M_k^t, \ell_k^t, q_k^t)\}_{k=1}^{K_t}$, where each mask M_k^t is paired with its matched prompt $\ell_k^t \in \mathcal{L}_t$ and confidence q_k^t . This ensures every mask carries a grounded label from the start, which our backend relies on for association and label assignment.

3.2 Backend for 3D Association and Track Representation

We adopt OpenVox’s [7] probabilistic voxel-consistency formulation as our association backbone, as voxel-level voting accumulates evidence across multiple observations, filtering out noisy detections over time, whereas greedy pairwise matching commits to early associations that can accumulate drift. The underlying voxel map is a sparse hash map that allocates cells only when RGB-D observations project points into them, so storage scales with observed surface rather than a prespecified volume. We adapt it for our open-vocabulary setting by replacing its SBERT caption pipeline with CLIP text-label embeddings for track association, computed directly from the detector’s label strings. These embeddings are invariant to viewpoint, unlike visual features. We therefore store visual embeddings separately for retrieval and grounding. Beyond the change of embedding, we introduce stricter observation filtering and confidence-weighted label voting to suppress noise and obtain more stable track representations.

Track scoring. Before association, we filter each 2D mask (erosion, subtraction of contained segments, area and confidence thresholds) and retain only the dominant 3D cluster via DBSCAN [8], which serves solely as a local per-mask denoising step. We then maintain a sparse voxel map that continuously records the spatial occupancy history of each track. Given a new observation $O = (M_k^t, \ell_k^t, q_k^t) \in \mathcal{O}_t$, we compute an association score $S(T_i, O)$ for each candidate track T_i using a weighted combination of geometric and feature similarities:

$$S(T_i, O) = \lambda_{\text{geo}} \cdot S_{\text{geo}} + \lambda_{\text{feat}} \cdot S_{\text{feat}}. \quad (1)$$

The geometric score S_{geo} represents a probabilistic voxel-consistency vote in the intersecting volume:

$$S_{\text{geo}}(T_i, O) = \frac{1}{|V(O)|} \sum_{v \in V(O)} \frac{c_{v,i}}{c_v}, \quad (2)$$

where $c_{v,i}$ is the number of times track T_i has been observed in voxel v , $c_v = \sum_j c_{v,j}$ is the total observation count across all tracks in that voxel, and $V(O)$ is the set of voxels occupied by observation O . The feature score S_{feat} is the cosine similarity between CLIP text embeddings of the incoming observation label and the track’s current consensus label. With $\lambda_{\text{geo}} = 0.8$ and $\lambda_{\text{feat}} = 0.2$, the observation is merged into the highest-scoring track if $S \geq \tau_{\text{assoc}}$. Otherwise, a new tentative track is instantiated. To prevent noisy detections from cluttering the scene graph, these tentative object tracks are promoted to confirmed nodes only after accumulating at least N_{conf} successful associations, effectively filtering out transient noise.

Confidence-weighted label voting. As observations accumulate, the same track may receive different label predictions across frames. To resolve these into a single consensus label, each track T_i maintains a label histogram $\mathcal{H}_i$ over the labels observed so far. When a new observation with grounded label ℓ and



Fig. 3: VLM Scheduling. Multi-target frame scheduling for the Description Agent (left) and pair scheduling for the Critic Agent (right), both using Set-of-Mark overlays.

detection confidence q is associated to T_i , the histogram is updated as

$$\mathcal{H}_i(\ell) \leftarrow \mathcal{H}_i(\ell) + q, \quad (3)$$

and the consensus label is the highest-scoring label:

$$\ell_i^* = \arg \max_{\ell} \mathcal{H}_i(\ell). \quad (4)$$

Weighting votes by detection confidence, rather than uniform counting, ensures that high-quality detections naturally suppress noisy or ambiguous predictions.

Visual embedding selection. To support downstream retrieval and grounding tasks, each track T_i aggregates a bank of CLIP visual embeddings $\mathcal{B}_i = \{\mathbf{f}_n^{\text{vis}}\}$, extracted from tight crops of its associated 2D segments. To ensure embedding quality, we apply two complementary filters. First, distant or heavily occluded viewpoints produce low-quality embeddings and should be discarded. An *adaptive area gate* maintains a running maximum segment area A_i^* over the track history and only admits semantic/CLIP updates whose visible 2D segment area A_n satisfies $A_n \geq \rho_{\text{area}} \cdot A_i^*$. Second, we select a single representative embedding via *text-guided scoring*. Relying solely on the largest crop is insufficient: For instance, a large view of a wall may still be dominated by occluding furniture, confounding the CLIP feature. Instead, we encode the track’s consensus label ℓ_i^* (Eq. (4)) with the CLIP text encoder to obtain a semantic anchor $\mathbf{f}_{\ell_i^*}^{\text{text}}$ and retrieve the most semantically aligned view:

$$\mathbf{f}_i^* = \mathbf{f}_{s^*}^{\text{vis}}, \text{ where } s^* = \arg \max_n \mathbf{f}_n^{\text{vis} \top} \mathbf{f}_{\ell_i^*}^{\text{text}}. \quad (5)$$

Spatial edges. To support relational grounding queries, we derive directed spatial edges deterministically from the axis-aligned bounding boxes of confirmed object tracks, following a similar approach to BBQ [19], as spatial relations can be reliably determined from geometric heuristics. Each edge encodes a directional predicate (e.g., left, under) and pairwise distance between target and anchor. Horizontal directions are resolved relative to a virtual room-center anchor to ensure viewpoint invariance.

3.3 Asynchronous Agentic Refinement

Enriching every node at every frame with a VLM would be prohibitively expensive, while deferring all enrichment to a post-hoc stage would mean that the system would not be queryable during exploration. We therefore introduce an asynchronous refinement layer comprising two VLM background agents that refine the graph without blocking the online loop (Fig. 3). The Critic Agent detects and merges duplicate object tracks caused by incremental association errors, while the Description Agent attaches fine-grained visual attributes to each node.

Critic agent. Similarly to pose drift in SLAM, incremental object associations accumulate errors over time, spawning duplicate or fragmented tracks. Prior consolidation methods [19] rely on feature similarity and 3D overlap, but embeddings struggle to distinguish distinct instances of the same class or to re-identify single objects across varying viewpoints. Resolving such an ambiguity requires both fine-grained visual evidence and a broader scene context, motivating the use of a VLM as a verifier. We therefore introduce the Critic Agent: an asynchronous in-the-loop verifier that performs merge-only identity corrections in two stages. First, a lightweight *candidate scheduling* step proposes high-confidence duplicate pairs. Then, a *visual verification* step authorizes or rejects each merge.

To avoid overwhelming the VLM with spurious proposals, we apply a two-gate filtering mechanism to each pair of tracks (T_i, T_j) , including tentative ones, with observation counts m_i, m_j (the number of successfully associated detections for each track) and 3D bounding boxes B_i, B_j . A *geometric gate* first discards any pair with $\min(m_i, m_j) < 2$ or $\text{IoU}_{3D}(B_i, B_j) < \tau_{\text{iou}}$, pruning noisy single-observation detections. Surviving pairs pass through a *semantic gate* that scores each pair by the cosine similarity s_{ij} between the SBERT [34] embeddings of their consensus labels. We use SBERT rather than CLIP here, as CLIP gives uniformly high similarity even for distinct labels, making threshold filtering unreliable.

Pairs with $s_{ij} \geq \tau_{\text{sem}}$ enter a candidate buffer that decouples tracking from verification. For each buffered pair, we track the peak alignment via a running maximum $\hat{s}_{ij} \leftarrow \max(\hat{s}_{ij}, s_{ij}^{(t)})$, capturing the best-observed compatibility as new evidence arrives. The buffer is flushed periodically, ranking candidates by $\hat{s}_{ij}$ and dispatching the top- K_{critic} to the verifier. At the end of the sequence, all remaining valid pairs are flushed to ensure complete coverage. When merging two tentative tracks results in a combined observation count that exceeds N_{conf} , the merged track is automatically promoted to a confirmed node.

For each selected pair, we first keep high-visibility views of each track, measured by visible mask area. If both tracks are visible in the same frame, we choose the shared frame where both are most visible. Otherwise, we choose two high-visibility views with similar camera direction, camera position, object depth, and temporal proximity. To ground VLM attention without obscuring texture details, we employ Set-of-Mark (SoM) prompting [50], overlaying high-contrast SAM mask boundaries and unique identifiers directly onto the image. The model receives these marked images alongside a context tuple $\mathcal{C} = \{\text{ID}_{i,j}, \text{Label}_{i,j}, \hat{s}_{ij}, \text{geometry cues}\}$ that supplies geometric cues to re-

duce hallucinations from ambiguous masks, and returns a structured decision $\mathcal{D} = \{a \in \{\text{MERGE}, \text{KEEP}\}, r \in \text{String}\}$. When the verifier returns $a = \text{MERGE}$, the Critic unifies point clouds, features, and label histograms of both tracks in the scene graph.

Description agent. While the tracker maintains object identities and coarse text labels, downstream tasks such as visual grounding require deeper semantic understanding. We introduce a Description Agent that runs asynchronously in a dedicated worker thread to enrich each scene graph node with fine-grained visual attributes, such as material, color, texture and physical state, which are critical for disambiguating instances of the same class.

Since a single frame typically depicts multiple objects, a small set of well-chosen frames can cover all targets while amortizing VLM cost and preserving scene context that individual crops would lack. We therefore introduce a multi-target frame scheduler that buffers 2D bounding boxes over a tumbling window of W_{desc} frames and identifies targets that lack a description or whose best visible area substantially exceeds their last-described area. It then greedily selects up to K_{desc} frames that maximize target coverage. A target T can be covered in frame f only if its normalized screen area $\hat{A}_{T,f}$ is within a fraction τ_{view} of its best view in the window. Each candidate frame is scored by

$$U(f) = \sum_{T \in \mathcal{T}_f} w_T \cdot \ln(1 + \gamma \hat{A}_{T,f}), \quad (6)$$

where $\mathcal{T}_f$ is the set of targets coverable in frame f , w_T is a per-target weight that down-weights background classes and γ controls the area sensitivity. The greedy loop selects the highest-scoring frame, marks its targets as covered, and repeats until the budget is exhausted.

Selected frames are also rendered with SoM overlays and submitted in a single multi-image VLM call, along with each target’s current label, description, and attribute estimates as context. For each marked object, the model returns a refined label (*e.g.*, “winged armchair” rather than “chair”) and visual attributes (material, color, state, texture). Attributes accumulate in a per-node histogram mirroring the confidence-weighted voting of Eq. (3), so repeated observations reinforce consistent evidence.

4 Experiments

We evaluate on open-vocabulary 3D semantic segmentation and 3D visual grounding, then analyze component contributions through ablations.

Datasets. For semantic segmentation, we evaluate on Replica [42] (eight photorealistic scenes: `room0–room2`, `office0–office4`) and ScanNet [6] (eight real-world RGB-D sequences: `0011_00`, `0030_00`, `0046_00`, `0086_00`, `0222_00`, `0378_00`, `0389_00`, `0435_00`). For visual grounding, we use three referring expression benchmarks with ScanNet: ScanRefer [5] (natural-language descriptions), Sr3D+ [1] (template-generated), and Nr3D [1] (free-form human-written). For both tasks,

we follow the evaluation protocol of BBQ [19]; the eight ScanNet scenes are chosen to match their grounding setup. We adopt the same frame sampling stride of 5 for Replica and 10 for ScanNet.

Metrics. For segmentation, we report mean per-class accuracy (mAcc), mean intersection-over-union (mIoU), and frequency-weighted IoU (f-mIoU), following the open-vocabulary protocol of BBQ [19] with EVA02-CLIP [9] features. Per-point predictions are matched to ground truth via nearest-neighbor association in 3D. Generic categories (*other*, *otherfurniture*) are excluded. For grounding, we use GPT-4o [27] and report accuracy at IoU thresholds A@0.25 and A@0.5 on ScanRefer, and at A@0.1 and A@0.25 on Sr3D+ and Nr3D. The full evaluator prompt and an ablation over evaluator models are provided in the supplementary.

Implementation. For our CLIP embeddings we use the EVA02-CLIP [9] variant. Qwen3-VL-2B-Instruct [2] runs locally; for the Agents we use GPT-5-mini [28], called via its cloud API. Experiments use a single workstation (AMD Ryzen 7 9700X, NVIDIA RTX 5090, 64 GB RAM). We set the voxel backend to a resolution of 4 cm with an adaptive area ratio $\rho_{\text{area}}=0.5$. Association balances geometry and feature similarity ($\lambda_{\text{geo}}=0.8$, $\lambda_{\text{feat}}=0.2$) at a threshold $\tau_{\text{assoc}}=0.4$; tentative tracks are promoted after $N_{\text{conf}}=6$ successful associations. The semantic gate is set to $\tau_{\text{sem}}=0.8$ and the IoU gate to $\tau_{\text{iou}}=0.01$, with utility scale $\gamma=100$. The Critic Agent flushes up to $K_{\text{critic}}=20$ candidate pairs every $W_{\text{critic}}=50$ frames. The Description Agent operates over tumbling windows of $W_{\text{desc}}=30$ frames, selecting up to $K_{\text{desc}}=3$ frames per window. The remaining hyperparameters are listed in the supplementary.

5 Results

We evaluate our method on two tasks: open-vocabulary 3D semantic segmentation and 3D visual grounding. We then analyze the contribution of individual components through ablations.

5.1 Open-Vocabulary 3D Semantic Segmentation

Tab. 1 compares our method with prior open-vocabulary 3D methods on Replica and ScanNet. On Replica, our incremental method outperforms all prior methods on every metric (0.58 mAcc, 0.37 mIoU, 0.61 f-mIoU), with the largest gain on f-mIoU (+0.13 over BBQ-CLIP). On ScanNet, we obtain 0.75 mAcc, 0.44 mIoU, and 0.46 f-mIoU, achieving the highest mIoU and f-mIoU while remaining competitive with OpenVox (0.76) on mAcc. We attribute these gains to two factors. First, the voxel-consistency association with CLIP text-label embeddings produces more stable

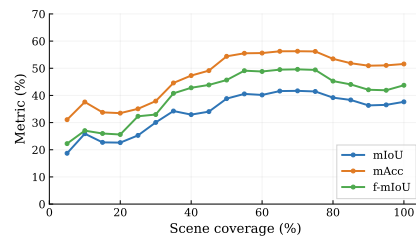


Fig. 4: Incremental Segmentation. On Room0 (Replica), metrics converge steadily as frames are processed.

Table 1: 3D Semantic Segmentation. Comparison with prior open-vocabulary methods on Replica and ScanNet.

Method	Replica			ScanNet		
	mAcc	mIoU	f-mIoU	mAcc	mIoU	f-mIoU
ConceptFusion [15]	0.29	0.11	0.14	0.49	0.26	0.31
OpenMask3D [43]	—	—	—	0.34	0.18	0.20
ConceptGraphs [11]	0.36	0.18	0.15	0.52	0.26	0.29
BBQ-CLIP [19]	0.38	0.27	0.48	0.56	0.34	0.36
Octree-Graph [44]	0.51	0.34	0.34	0.71	0.40	0.37
OpenVox [7]	0.51	0.26	0.40	0.76	0.43	0.39
ThinkGraphs	0.58	0.37	0.61	0.75	0.44	0.46

track identities than the pairwise visual similarity used by prior methods. Second, the text-guided feature selection in Eq. (5) ensures that each track is represented by the view most consistent with its consensus label, producing cleaner per-point semantic assignments. Fig. 4 illustrates this effect over time: Metrics improve steadily as additional frames are processed, demonstrating that the incremental pipeline converges to high-quality predictions without requiring a complete reconstruction.

5.2 3D Visual Grounding

We evaluate on three referring expression benchmarks of increasing difficulty: the template-based Sr3D+, the free-form Nr3D, and the natural-language ScanRefer. To handle the complexity of free-form and natural-language evaluation, we utilize GPT-4o [27] as our automated evaluator.

Sr3D+ and Nr3D. On both referral benchmarks (Tab. 2), our method substantially outperforms all baselines, improving over Open3DSG [17] by 15.3 A@0.25 on Sr3D+ and 18.8 on Nr3D. Performance remains consistent across difficulty and viewpoint splits (see supplementary for per-split breakdowns).

ScanRefer. On the ScanRefer benchmark (Tab. 3), our method achieves 52.9% A@0.25 and 37.6% A@0.5, outperforming Open3DSG [17] by 18.3 and 6.6 absolute points respectively. This reflects the combined effect of our core components: our backend provides stable track identities and high-quality representative visual embeddings, the Critic Agent reduces node fragmentation so that queries match the correct consolidated object, and the Description Agent provides fine-grained attributes that help disambiguate same-class instances.

5.3 Ablation Studies

To isolate individual component contributions, we conduct ablations on the eight Replica scenes.

Table 2: 3D Visual Grounding. Overall accuracy on Sr3D+ and Nr3D.

Method	Sr3D+		Nr3D	
	A@0.1	A@0.25	A@0.1	A@0.25
OpenFusion [48]	12.6	2.4	10.7	1.4
BBQ-CLIP [19]	14.4	8.8	15.3	9.4
ConceptGraphs [11]	13.3	6.2	16.0	7.2
BBQ [19]	34.2	22.7	28.3	19.0
Open3DSG [17]	28.9	27.7	25.9	23.0
ThinkGraphs	48.6	43.0	47.4	41.8

Table 3: 3D Visual grounding. Overall accuracy on ScanRefer.

Method	ScanRefer	
	A@0.25	A@0.5
LERF [16]	4.4	0.3
OpenScene [30]	13.0	5.1
LLM-Grounder [49]	17.1	5.3
BBQ [19]	19.4	11.6
Open3DSG [17]	34.6	31.0
ThinkGraphs	52.9	37.6

Table 4: Pipeline ablation. Component contributions on Replica; after replacing the RAM++ baseline with Qwen, components are added incrementally.

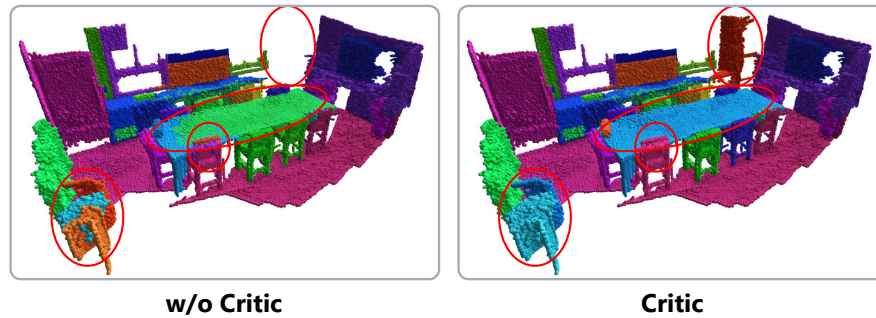
Configuration	mAcc	mIoU	f-mIoU
RAM++ (baseline frontend)	0.37	0.20	0.34
Qwen (our frontend)	0.46	0.25	0.43
+ CLIP-text labels	0.47	0.26	0.44
+ Text-guided selection	0.52	0.35	0.55
+ Area filter	0.58	0.37	0.61

Pipeline design choices. Tab. 4 ablates the backend pipeline on Replica. Switching from RAM++ [13] to Qwen3-VL [2] as the frontend tagger yields the largest mAcc gain (+0.09), as context-aware prompts produce fewer missed detections. Text-guided embedding selection drives the largest mIoU gain (+0.09): Naively selecting the largest-area crop as the representative embedding, as done *e.g.* in BBQ [19], is sensitive to foreground clutter. Scoring against the consensus label produces cleaner views. The adaptive area gate adds a further 0.06 mAcc by filtering low-quality observations, while replacing SBERT with CLIP text-label embeddings provides a smaller benefit but unifies the visual and textual feature space.

Agent contributions to grounding. Tab. 5 incrementally adds the two asynchronous agents to isolate their effect on Nr3D grounding accuracy. The base system includes the full backend with spatial relations, but no VLM agents. Adding the Critic Agent alone lifts overall A@0.25 from 38.5% to 39.3%, with gains concentrated on easy and view-independent splits where fragmented object tracks are the primary bottleneck (Fig. 5). The Description Agent has a larger effect on hard referrals (+3.5 A@0.25), where fine-grained attributes are needed to disambiguate same-class instances. Combining both agents yields the best results in all splits (41.8% A@0.25 overall, +7.1 on hard), confirming that the two agents address complementary failure modes. An ablation isolating the

Table 5: Agent ablation. Effect of the Critic and Description agents on Nr3D; Base is the online mapping pipeline without the VLM agents.

Method	Overall		Easy		Hard		View Dep.		View Indep.	
	A@0.1	A@0.25	A@0.1	A@0.25	A@0.1	A@0.25	A@0.1	A@0.25	A@0.1	A@0.25
Base	43.1	38.5	49.2	46.2	27.4	18.8	43.0	37.6	43.1	38.8
Base + Critic	45.1	39.3	52.2	47.4	26.9	18.8	43.0	38.2	45.7	39.7
Base + Description	44.5	40.5	50.0	47.6	30.5	22.3	44.8	40.6	44.4	40.4
Base + Critic + Description	47.4	41.8	53.0	48.0	33.0	25.9	49.1	43.0	46.8	41.4

**Fig. 5: Effect of the Critic Agent.** Without the Critic (left), association drift fragments objects into duplicate object tracks (red circles). Semantic loop closure (right) merges them into consistent identities. Points are colored by instance.

architecture from the frontend VLM (a lighter RAM++ tagger) is provided in the supplementary, showing the gains are not merely due to a stronger tagger.

Runtime and cost. The online mapping path averages 1.45s per keyframe, with the Critic and Description agents running off the critical path so they never block the online mapping loop. The multi-target scheduler keeps VLM usage low, about 26 calls per scene on average, more than two orders of magnitude fewer than naive per-object-per-keyframe prompting ($\sim 9,300$ on a typical scene). The primary runtime bottleneck is the frontend, not the agents. A full analysis is provided in the supplementary.

5.4 Limitations

Our framework has two main limitations. Firstly, the front-end models (Qwen3-VL + Grounded-SAM) prevent real-time operation. Replacing these models with more efficient alternatives would be the most direct way of overcoming this issue. Secondly, asynchronous agents inevitably lag behind the mapping frontier, meaning that objects discovered during fast exploration may remain unenriched for several frames. This limits the quality of grounding early in a sequence.

6 Conclusion

We present ThinkGraphs, an asynchronous architecture for open-vocabulary 3D scene graphs that decouples lightweight online mapping from heavyweight VLM reasoning. By running semantic refinement in the background, the scene graph remains queryable during exploration and grows in richness over time. A probabilistic voxel backbone with text-guided embedding selection provides stable object identities, while a Critic Agent corrects fragmentation through semantic loop closure, and a Description Agent attaches fine-grained attributes via multi-target frame scheduling. Our method outperforms all batch-based open-vocabulary scene-graph methods on semantic segmentation on Replica and ScanNet, and surpasses the prior state-of-the-art across three visual grounding benchmarks (Sr3D+, Nr3D, ScanRefer) by 15.3 to 18.8 A@0.25, showing that an incremental architecture can match and exceed batch-based alternatives. Future work could extend this asynchronous agent design with additional specialized agents, for example for richer relation prediction, higher-level scene reasoning, or task-oriented graph refinement, while also exploring its use in other structured representations that require both responsiveness and semantic depth.

Acknowledgements

The authors thank the International Max Planck Research School for Intelligent Systems (IMPRS-IS) for supporting Deniz Bickici and Michael Pabst. This work was supported by the Alexander von Humboldt Foundation funded by the German Federal Ministry of Research, Technology and Space, and by Deutsche Forschungsgemeinschaft (DFG) under Germany’s Excellence Strategy EXC 2120/2 (390831618).

References

1. Achlioptas, P., Abdelreheem, A., Xia, F., Elhoseiny, M., Guibas, L.: ReferIt3D: Neural Listeners for Fine-Grained 3D Object Identification in Real-World Scenes. In: Computer Vision – ECCV 2020, vol. 12346, pp. 422–440. Springer International Publishing, Cham (2020). https://doi.org/10.1007/978-3-030-58452-8_25
2. Bai, S., Cai, Y., Chen, R., Chen, K., Chen, X., Cheng, Z., Deng, L., Ding, W., Gao, C., Ge, C., Ge, W., Guo, Z., Huang, Q., Huang, J., Huang, F., Hui, B., Jiang, S., Li, Z., Li, M., Li, M., Li, K., Lin, Z., Lin, J., Liu, X., Liu, J., Liu, C., Liu, Y., Liu, D., Liu, S., Lu, D., Luo, R., Lv, C., Men, R., Meng, L., Ren, X., Ren, X., Song, S., Sun, Y., Tang, J., Tu, J., Wan, J., Wang, P., Wang, P., Wang, Q., Wang, Y., Xie, T., Xu, Y., Xu, H., Xu, J., Yang, Z., Yang, M., Yang, J., Yang, A., Yu, B., Zhang, F., Zhang, H., Zhang, X., Zheng, B., Zhong, H., Zhou, J., Zhou, F., Zhou, J., Zhu, Y., Zhu, K.: Qwen3-VL Technical Report (2025). <https://doi.org/10.48550/arXiv.2511.21631>
3. Bavle, H., Sanchez-Lopez, J.L., Shaheer, M., Civera, J., Voos, H.: *S-Graphs+*: Real-Time Localization and Mapping Leveraging Hierarchical Representations. IEEE Robotics and Automation Letters **8**(8), 4927–4934 (2023). <https://doi.org/10.1109/LRA.2023.3290512>

4. Chen, B., Xia, F., Ichter, B., Rao, K., Gopalakrishnan, K., Ryoo, M., Stone, A., Kappler, D.: Open-vocabulary Queryable Scene Representations for Real World Planning. In: 2023 IEEE International Conference on Robotics and Automation (ICRA). pp. 11509–11522. IEEE, London, United Kingdom (2023). <https://doi.org/10.1109/ICRA48891.2023.10161534>
5. Chen, D.Z., Chang, A.X., Nießner, M.: ScanRefer: 3D Object Localization in RGB-D Scans Using Natural Language. In: Computer Vision – ECCV 2020, vol. 12365, pp. 202–221. Springer International Publishing, Cham (2020). https://doi.org/10.1007/978-3-030-58565-5_13
6. Dai, A., Chang, A.X., Savva, M., Halber, M., Funkhouser, T., Nießner, M.: ScanNet: Richly-Annotated 3D Reconstructions of Indoor Scenes. In: 2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR). pp. 2432–2443. IEEE, Honolulu, HI (2017). <https://doi.org/10.1109/CVPR.2017.261>
7. Deng, Y., Yao, B., Tang, Y., Zhou, T., Yang, Y., Yue, Y.: OpenVox: Real-time instance-level open-vocabulary probabilistic voxel representation. In: 2025 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS). pp. 1305–1311 (2025). <https://doi.org/10.1109/IROS60139.2025.11246455>
8. Ester, M., Kriegel, H.P., Sander, J., Xu, X.: A density-based algorithm for discovering clusters in large spatial databases with noise. In: Proceedings of the Second International Conference on Knowledge Discovery and Data Mining (KDD). pp. 226–231. AAAI Press, Portland, Oregon (1996)
9. Fang, Y., Sun, Q., Wang, X., Huang, T., Wang, X., Cao, Y.: Eva-02: A visual representation for neon genesis. *Image Vision Comput.* **149**(C) (2024). <https://doi.org/10.1016/j.imavis.2024.105171>
10. Feng, M., Yan, C., Wu, Z., Dong, W., Wang, Y., Mian, A.: History-enhanced 3d scene graph reasoning from rgb-d sequences. *IEEE Transactions on Circuits and Systems for Video Technology* **35**(8), 7667–7682 (2025). <https://doi.org/10.1109/TCSVT.2025.3548308>
11. Gu, Q., Kuwajerwala, A., Morin, S., Jatavallabhula, K.M., Sen, B., Agarwal, A., Rivera, C., Paul, W., Ellis, K., Chellappa, R., Gan, C., De Melo, C.M., Tenenbaum, J.B., Torralba, A., Shkurti, F., Paull, L.: ConceptGraphs: Open-Vocabulary 3D Scene Graphs for Perception and Planning. In: 2024 IEEE International Conference on Robotics and Automation (ICRA). pp. 5021–5028. IEEE, Yokohama, Japan (2024). <https://doi.org/10.1109/ICRA57147.2024.10610243>
12. Huang, C., Mees, O., Zeng, A., Burgard, W.: Visual Language Maps for Robot Navigation. In: 2023 IEEE International Conference on Robotics and Automation (ICRA). pp. 10608–10615 (2023). <https://doi.org/10.1109/ICRA48891.2023.10160969>
13. Huang, X., Huang, Y.J., Zhang, Y., Tian, W., Feng, R., Zhang, Y., Xie, Y., Li, Y., Zhang, L.: Open-set image tagging with multi-grained text supervision. In: Proceedings of the 33rd ACM International Conference on Multimedia (ACM MM). p. 4117–4126. Association for Computing Machinery, New York, NY, USA (2025). <https://doi.org/10.1145/3746027.3755316>
14. Hughes, N., Chang, Y., Carlone, L.: Hydra: A Real-time Spatial Perception System for 3D Scene Graph Construction and Optimization. In: Robotics: Science and Systems XVIII. Robotics: Science and Systems Foundation (2022). <https://doi.org/10.15607/RSS.2022.XVIII.050>
15. Jatavallabhula, K., Kuwajerwala, A., Gu, Q., Omama, M., Iyer, G., Saryazdi, S., Chen, T., Maalouf, A., Li, S., Keetha, N., Tewari, A., Tenenbaum, J., Melo, C., Krishna, M., Paull, L., Shkurti, F., Torralba, A.: ConceptFusion: Open-set multi-

- modal 3D mapping. In: *Robotics: Science and Systems XIX. Robotics: Science and Systems Foundation* (2023). <https://doi.org/10.15607/RSS.2023.XIX.066>
16. Kerr, J., Kim, C.M., Goldberg, K., Kanazawa, A., Tancik, M.: LERF: Language Embedded Radiance Fields. In: *2023 IEEE/CVF International Conference on Computer Vision (ICCV)*. pp. 19672–19682. IEEE, Paris, France (2023). <https://doi.org/10.1109/ICCV51070.2023.01807>
 17. Koch, S., Vaskevicius, N., Colosi, M., Hermosilla, P., Ropinski, T.: Open3dsg: Open-vocabulary 3d scene graphs from point clouds with queryable objects and open-set relationships. In: *2024 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. pp. 14183–14193 (2024). <https://doi.org/10.1109/CVPR52733.2024.01345>
 18. Lin, S., Wang, J., Xu, M., Zhao, H., Chen, Z.: Topology Aware Object-Level Semantic Mapping Towards More Robust Loop Closure. *IEEE Robotics and Automation Letters* **6**(4), 7041–7048 (2021). <https://doi.org/10.1109/LRA.2021.3097242>
 19. Linok, S., Zemskova, T., Ladanova, S., Titkov, R., Yudin, D., Monastyrny, M., Valenkov, A.: Beyond Bare Queries: Open-Vocabulary Object Grounding with 3D Scene Graph. In: *2025 IEEE International Conference on Robotics and Automation (ICRA)*. pp. 13582–13589 (2025). <https://doi.org/10.1109/ICRA55743.2025.11128059>
 20. Liu, P., Guo, Z., Warke, M., Chintala, S., Paxton, C., Shafiullah, N.M.M., Pinto, L.: Dynamem: Online Dynamic Spatio-Semantic Memory for Open World Mobile Manipulation. In: *2025 IEEE International Conference on Robotics and Automation (ICRA)*. pp. 13346–13355 (2025). <https://doi.org/10.1109/ICRA55743.2025.11127619>
 21. Liu, S., Zeng, Z., Ren, T., Li, F., Zhang, H., Yang, J., Jiang, Q., Li, C., Yang, J., Su, H., Zhu, J., Zhang, L.: Grounding DINO: Marrying DINO with Grounded Pre-training for Open-Set Object Detection. In: *Computer Vision – ECCV 2024*, vol. 15105, pp. 38–55. Springer Nature Switzerland, Cham (2025). https://doi.org/10.1007/978-3-031-72970-6_3
 22. Liu, Y., Petillot, Y., Lane, D., Wang, S.: Global Localization with Object-Level Semantics and Topology. In: *2019 International Conference on Robotics and Automation (ICRA)*. pp. 4909–4915. IEEE, Montreal, QC, Canada (2019). <https://doi.org/10.1109/ICRA.2019.8794475>
 23. Ma, X., Yong, S., Zheng, Z., Li, Q., Liang, Y., Zhu, S.C., Huang, S.: Sqa3d: Situated question answering in 3d scenes. In: *International Conference on Learning Representations (ICLR)* (2023)
 24. Maggio, D., Chang, Y., Hughes, N., Trang, M., Griffith, D., Dougherty, C., Cristofalo, E., Schmid, L., Carlone, L.: Clio: Real-Time Task-Driven Open-Set 3D Scene Graphs. *IEEE Robotics and Automation Letters* **9**(10), 8921–8928 (2024). <https://doi.org/10.1109/LRA.2024.3451395>
 25. McCormac, J., Clark, R., Bloesch, M., Davison, A., Leutenegger, S.: Fusion++: Volumetric object-level slam. In: *2018 International Conference on 3D Vision (3DV)*. pp. 32–41 (2018). <https://doi.org/10.1109/3DV.2018.00015>
 26. Mei, G., Riz, L., Wang, Y., Poiesi, F.: Vocabulary-Free 3D Instance Segmentation with Vision-Language Assistant. In: *2025 International Conference on 3D Vision (3DV)*. pp. 1197–1210 (2025). <https://doi.org/10.1109/3DV66043.2025.00114>
 27. OpenAI: GPT-4o System Card (2024). <https://doi.org/10.48550/arXiv.2410.21276>
 28. OpenAI: OpenAI GPT-5 System Card (2025). <https://doi.org/10.48550/arXiv.2601.03267>

29. Oquab, M., Darcet, T., Moutakanni, T., Vo, H.V., Szafraniec, M., Khalidov, V., Fernandez, P., Haziza, D., Massa, F., El-Nouby, A., Assran, M., Ballas, N., Galuba, W., Howes, R., Huang, P.Y., Li, S.W., Misra, I., Rabbat, M., Sharma, V., Synnaeve, G., Xu, H., Jegou, H., Mairal, J., Labatut, P., Joulin, A., Bojanowski, P.: DINOv2: Learning robust visual features without supervision. *Transactions on Machine Learning Research* (2024)
30. Peng, S., Genova, K., Jiang, C., Tagliasacchi, A., Pollefeys, M., Funkhouser, T.: OpenScene: 3D Scene Understanding with Open Vocabularies. In: 2023 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 815–824. IEEE, Vancouver, BC, Canada (2023). <https://doi.org/10.1109/CVPR52729.2023.00085>
31. Peterson, M.B., Jia, Y.X., Tian, Y., Thomas, A., How, J.P.: Roman: Open-set object map alignment for robust view-invariant global localization. In: *Robotics: Science and Systems (RSS)* (2025). <https://doi.org/10.15607/RSS.2025.XXI.029>
32. Radford, A., Kim, J.W., Hallacy, C., Ramesh, A., Goh, G., Agarwal, S., Sastry, G., Askell, A., Mishkin, P., Clark, J., Krueger, G., Sutskever, I.: Learning transferable visual models from natural language supervision. In: *Proceedings of the 38th International Conference on Machine Learning (ICML)*. vol. 139, pp. 8748–8763. PMLR (2021)
33. Ravi, N., Gabeur, V., Hu, Y.T., Hu, R., Ryali, C., Ma, T., Khedr, H., Rädle, R., Rolland, C., Gustafson, L., Mintun, E., Pan, J., Alwala, K.V., Carion, N., Wu, C.Y., Girshick, R., Dollár, P., Feichtenhofer, C.: SAM 2: Segment Anything in Images and Videos (2024). <https://doi.org/10.48550/ARXIV.2408.00714>
34. Reimers, N., Gurevych, I.: Sentence-BERT: Sentence embeddings using Siamese BERT-networks. In: *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*. pp. 3982–3992. Association for Computational Linguistics, Hong Kong, China (2019). <https://doi.org/10.18653/v1/D19-1410>
35. Ren, T., Liu, S., Zeng, A., Lin, J., Li, K., Cao, H., Chen, J., Huang, X., Chen, Y., Yan, F., Zeng, Z., Zhang, H., Li, F., Yang, J., Li, H., Jiang, Q., Zhang, L.: Grounded SAM: Assembling Open-World Models for Diverse Visual Tasks (2024). <https://doi.org/10.48550/arXiv.2401.14159>
36. Renz, M., Igelbrink, F., Atzmueller, M.: Integrating prior observations for incremental 3d scene graph prediction. In: 2025 International Conference on Machine Learning and Applications (ICMLA). pp. 887–892 (2025). <https://doi.org/10.1109/ICMLA66185.2025.00132>
37. Rosinol, A., Abate, M., Chang, Y., Carlone, L.: Kimera: an Open-Source Library for Real-Time Metric-Semantic Localization and Mapping. In: 2020 IEEE International Conference on Robotics and Automation (ICRA). pp. 1689–1696 (2020). <https://doi.org/10.1109/ICRA40945.2020.9196885>
38. Rosinol, A., Violette, A., Abate, M., Hughes, N., Chang, Y., Shi, J., Gupta, A., Carlone, L.: Kimera: From SLAM to spatial perception with 3D dynamic scene graphs. *The International Journal of Robotics Research* **40**(12-14), 1510–1546 (2021). <https://doi.org/10.1177/02783649211056674>
39. Runz, M., Agapito, L.: Co-fusion: Real-time segmentation, tracking and fusion of multiple objects. In: 2017 IEEE International Conference on Robotics and Automation (ICRA). pp. 4471–4478. IEEE, Singapore, Singapore (2017). <https://doi.org/10.1109/ICRA.2017.7989518>

40. Runz, M., Buffier, M., Agapito, L.: MaskFusion: Real-Time Recognition, Tracking and Reconstruction of Multiple Moving Objects. In: 2018 IEEE International Symposium on Mixed and Augmented Reality (ISMAR). pp. 10–20. IEEE, Munich, Germany (2018). <https://doi.org/10.1109/ISMAR.2018.00024>
41. Shafiullah, N.M.M., Paxton, C., Pinto, L., Chintala, S., Szlam, A.: CLIP-Fields: Weakly Supervised Semantic Fields for Robotic Memory. In: Robotics: Science and Systems XIX (2023). <https://doi.org/10.15607/RSS.2023.XIX.074>
42. Straub, J., Whelan, T., Ma, L., Chen, Y., Wijmans, E., Green, S., Engel, J.J., Mur-Artal, R., Ren, C., Verma, S., Clarkson, A., Yan, M., Budge, B., Yan, Y., Pan, X., Yon, J., Zou, Y., Leon, K., Carter, N., Briales, J., Gillingham, T., Mueggler, E., Pesqueira, L., Savva, M., Batra, D., Strasdat, H.M., Nardi, R.D., Goesele, M., Lovegrove, S., Newcombe, R.: The Replica Dataset: A Digital Replica of Indoor Spaces (2019). <https://doi.org/10.48550/arXiv.1906.05797>
43. Takmaz, A., Fedele, E., Sumner, R.W., Pollefeys, M., Tombari, F., Engelmann, F.: Openmask3d: open-vocabulary 3d instance segmentation. In: Proceedings of the 37th International Conference on Neural Information Processing Systems (NeurIPS). Curran Associates Inc., Red Hook, NY, USA (2023)
44. Wang, Z., Su, Y., Li, C., Wang, D., Huang, Y., Li, X., Zhao, B.: Open-vocabulary octree-graph for 3d scene understanding. In: 2025 IEEE/CVF International Conference on Computer Vision (ICCV). pp. 7037–7047 (2025). <https://doi.org/10.1109/ICCV51701.2025.00661>
45. Werby, A., Huang, C., Büchner, M., Valada, A., Burgard, W.: Hierarchical Open-Vocabulary 3D Scene Graphs for Language-Grounded Robot Navigation. In: Robotics: Science and Systems XX. Robotics: Science and Systems Foundation (2024). <https://doi.org/10.15607/RSS.2024.XX.077>
46. Wu, S.C., Tateno, K., Navab, N., Tombari, F.: Incremental 3D Semantic Scene Graph Prediction from RGB Sequences. In: 2023 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 5064–5074. IEEE, Vancouver, BC, Canada (2023). <https://doi.org/10.1109/CVPR52729.2023.00490>
47. Wu, S.C., Wald, J., Tateno, K., Navab, N., Tombari, F.: SceneGraphFusion: Incremental 3D Scene Graph Prediction from RGB-D Sequences. In: 2021 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 7511–7521. IEEE, Nashville, TN, USA (2021). <https://doi.org/10.1109/CVPR46437.2021.00743>
48. Yamazaki, K., Hanyu, T., Vo, K., Pham, T., Tran, M., Doretto, G., Nguyen, A., Le, N.: Open-Fusion: Real-time Open-Vocabulary 3D Mapping and Queryable Scene Representation. In: 2024 IEEE International Conference on Robotics and Automation (ICRA). pp. 9411–9417 (2024). <https://doi.org/10.1109/ICRA57147.2024.10610193>
49. Yang, J., Chen, X., Qian, S., Madaan, N., Iyengar, M., Fouhey, D.F., Chai, J.: LLM-Grounder: Open-Vocabulary 3D Visual Grounding with Large Language Model as an Agent. In: 2024 IEEE International Conference on Robotics and Automation (ICRA). pp. 7694–7701 (2024). <https://doi.org/10.1109/ICRA57147.2024.10610443>
50. Yang, J., Zhang, H., Li, F., Zou, X., Li, C., Gao, J.: Set-of-Mark Prompting Unleashes Extraordinary Visual Grounding in GPT-4V (2023). <https://doi.org/10.48550/arXiv.2310.11441>
51. Zhang, C., Delitzas, A., Wang, F., Zhang, R., Ji, X., Pollefeys, M., Engelmann, F.: Open-Vocabulary Functional 3D Scene Graphs for Real-World Indoor Spaces. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) (2025). <https://doi.org/10.1109/CVPR52734.2025.01807>