



EgoNight: Towards Egocentric Vision Understanding at Night with a Challenging Benchmark

Deheng Zhang^{1*}, Yuqian Fu^{1*†}, Runyi Yang¹, Yang Miao¹, Tianwen Qian², Xu Zheng^{1,3}, Guolei Sun⁴,
Ajad Chhatkuli¹, Xuanjing Huang⁵, Yu-Gang Jiang⁵, Luc Van Gool¹, Danda Pani Paudel¹

¹INSAIT, Sofia University “St. Kliment Ohridski”

²East China Normal University

³HKUST(GZ)

⁴Nankai University

⁵Fudan University

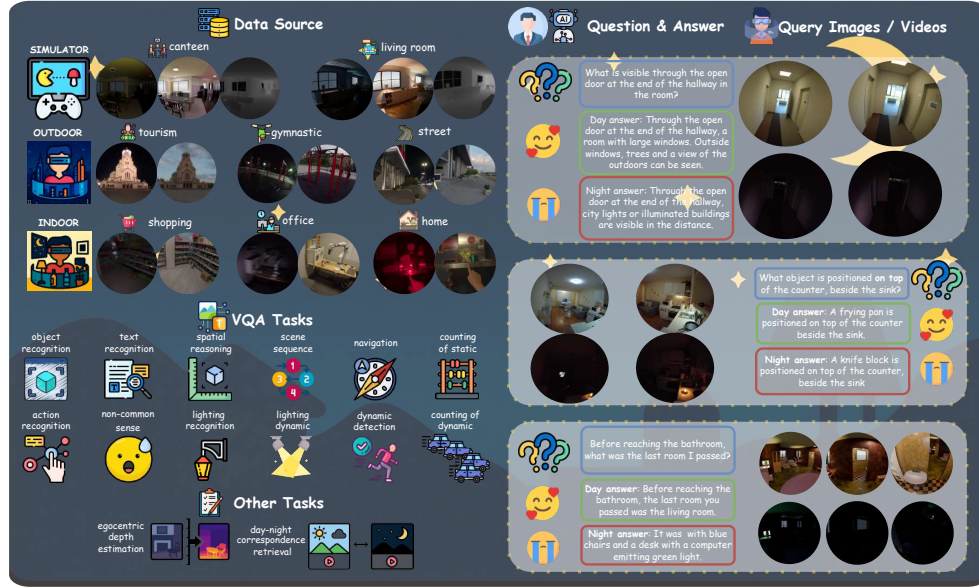


Figure 1: **Overview of the EgoNight.** EgoNight integrates diverse video sources spanning synthetic environments, real-world indoor and outdoor scenes, recorded under both daytime and nighttime conditions, with spatial and temporal alignment. It consists of three benchmarks: (i) *egocentric VQA* as the primary focus, (ii) *day–night correspondence retrieval*, and (iii) *egocentric depth estimation*, all targeting the challenges of low-light egocentric vision. The day–night alignment (illustrated on the right with VQA examples) enables rigorous analysis of illumination gaps in MLLMs.

Abstract

Most existing benchmarks for egocentric vision understanding focus primarily on daytime scenarios, overlooking the low-light conditions that are inevitable in real-world applications. To investigate this gap, we present **EgoNight**, the first comprehensive benchmark for nighttime egocentric vision, with visual question answering (VQA) as the core task. A key feature of EgoNight is the introduction of day–night aligned videos, which enhance night annotation quality using the daytime data and reveal clear performance gaps between lighting conditions. To achieve this, we collect both synthetic videos rendered by Blender and real-world recordings, ensuring that scenes and actions are visually and temporally aligned. Leveraging these paired videos, we construct **EgoNight-VQA**, supported by a novel

* means equal contribution; † denotes corresponding author.

day-augmented night auto-labeling engine and refinement through extensive human verification. Each QA pair is double-checked by annotators for reliability. In total, EgoNight-VQA contains 3658 QA pairs across 90 videos, spanning 12 diverse QA types, with more than 300 hours of human work. Evaluations of the state-of-the-art multimodal large language models (MLLMs) reveal substantial performance drops when transferring from day to night, underscoring the challenges of reasoning under low-light conditions. Beyond VQA, EgoNight also introduces two auxiliary tasks, *day–night correspondence retrieval* and *egocentric depth estimation at night*, that further explore the boundaries of existing models. We believe EgoNight-VQA provides a strong foundation for advancing application-driven egocentric vision research and for developing models that generalize across illumination domains. All the data and code will be made available upon acceptance.

1 Introduction

With the rapid development of wearable devices, egocentric vision understanding has become increasingly important. Unlike third-person vision, egocentric perception naturally aligns with the way humans perceive, understand, and interact with the world. A robust egocentric vision system can not only serve as an intelligent assistant in daily activities [1] but also play a crucial role in embodied AI such as robotics [2–4] and autonomous driving [5]. Beyond these general applications, egocentric vision holds unique potential for assisting specific user groups such as people who are blind or visually impaired [6], or physically disabled [7], enabling technologies that enhance navigation, accessibility, and real-time scene understanding.

Significant efforts have been made to advance egocentric vision understanding, including the construction of large-scale ego-centric datasets such as EPIC-KITCHENS [8], Ego4D [9], and Ego-Exo4D [10]; the design of diverse and challenging benchmarks such as EgoTaskQA [11], EgoSchema [12], EgoTempo [13], and EgoCross [14]; and the development of egocentric multimodal large language models (MLLMs) such as EgoVLPv2 [15], EgoGPT [1], and Exo2Ego [16]. Despite these advances, almost all prior works focus on daytime scenarios with favorable lighting. In contrast, real-world egocentric systems, for example, intelligent personal assistants for navigation, must operate at night, under low light, uneven illumination, and severely limited visibility. This motivates us to *investigate egocentric vision at night, with an emphasis on complex scene understanding and reasoning tasks*.

A central challenge in constructing such a benchmark lies in obtaining suitable video sources that capture the characteristics of nighttime environments, as well as developing annotation methods that ensure high labeling quality. To address this, we place particular emphasis on *day–night aligned videos*, which not only allow us to leverage daytime data to annotate nighttime videos, but also enable rigorous performance comparisons across day and night lighting conditions. However, in practice, collecting perfectly aligned day–night pairs in the real world is highly nontrivial. To overcome this, we leverage Blender [17], where scene layouts, camera trajectories, and lighting can be precisely controlled, enabling the synthesis of the desired videos. This produces **EgoNight-Synthetic**, a collection of 50 ideally aligned egocentric pairs spanning diverse and complex indoor scenarios with varying illumination levels. To complement synthetic data with real-world evidence, we design a *video-guided recording protocol* to construct **EgoNight-Sofia**, which contains 20 pairs of real-world egocentric videos with spatially and temporally aligned day–night counterparts. These videos cover realistic use cases (e.g., “Where did I put my keys?”, “How much is the item I saw in the grocery shop?”), spanning both indoor and outdoor environments under diverse illumination sources such as streetlights, flashlights, and candles. Finally, we incorporate 20 nighttime videos from the Oxford Day-and-Night dataset [18], termed **EgoNight-Oxford**, which serve as an additional testbed despite lacking day–night alignment. Together, these three video sources constitute our **EgoNight** dataset, which is the first egocentric dataset providing day–night aligned correspondences, as mainly summarized in Fig. 1.

The videos in EgoNight pave the way for constructing challenging benchmarks to evaluate the capabilities of existing models. Among many egocentric tasks, we focus on the egocentric video question answering, a flagship task that best reflects high-level understanding in egocentric vision. Specifically, to comprehensively evaluate model abilities, we first propose a diverse set of QA types, spanning well-studied tasks (e.g., object recognition, spatial reasoning, action recognition,

counting, text recognition) as well as several underexplored dimensions (e.g., temporal scene sequence understanding, navigation, lighting recognition, and non-common-sense reasoning). These are further organized into paired and unpaired QA types, depending on whether day-night counterparts share the same questions and answers. To construct the benchmark at scale, we then develop a *novel three-stage day-augmented auto-labeling pipeline* that leverages daytime videos to assist in generating question-answer pairs for nighttime clips, followed by extensive human verification to ensure accuracy and reliability. Building EgoNight and annotating VQA required **over 300 hours of human effort**, with each QA pair verified by at least one expert annotator. This process results in the high-quality **EgoNight-VQA** dataset, comprising 3,658 QA pairs. Beyond VQA, we introduce two auxiliary tasks with dedicated testbeds: day-night correspondence retrieval, which evaluates cross-illumination matching, and egocentric depth estimation at night, which is crucial for navigation and interaction in embodied AI. These two tasks further broaden the benchmark and expose new challenges for existing models.

Our extensive experiments across three video sources, three tasks, and 10 state-of-the-art multimodal large language models reveal that nearly all models (including closed-source models such as GPT and Gemini) struggle on this challenging benchmark, with a clear and consistent performance gap between day and night. This highlights the unsolved challenges of egocentric vision at nighttime and calls for more robust models that generalize across illumination conditions. Besides, we highlight that our newly introduced QA types, covering lighting recognition/dynamic, scene sequence reasoning, navigation, and non-common-sense reasoning, are substantially more challenging than well-studied categories, revealing fresh difficulties for MLLMs.

Our main contributions are threefold: i) **EgoNight Dataset**: We present the first egocentric dataset that systematically addresses nighttime conditions, featuring day-night aligned videos from synthetic (EgoNight-Synthetic), real-world (EgoNight-Sofia), and existing (EgoNight-Oxford) sources. ii) **Benchmark Suite**: We build a comprehensive benchmark centered on egocentric VQA with diverse QA types and 3658 fully human-verified QA pairs, complemented by egocentric depth estimation at night and day-night correspondence retrieval tasks. iii) **Empirical Insights**: Extensive evaluations reveal clear day-night performance gaps, underscoring illumination robustness as a key challenge; our newly proposed QA types are also validated to pose practical difficulties for current MLLMs.

2 Related Works

2.1 Egocentric Datasets and VQA Benchmarks

A series of large-scale egocentric datasets, such as EPIC-KITCHENS [8], Ego4D [9], Ego-Exo4D [10], and EgoExoLearn [19], have laid the foundation for a wide range of tasks, including action recognition [20], object detection [21], pose estimation [22], video generation [23], Ego-Exo correspondence [24]. Among these, we are particularly interested in egocentric visual question answering (VQA) [25], which provides a natural and human-like framework for comprehensively evaluating model performance through question-answer interactions. In recent years, several egocentric VQA benchmarks have been proposed, including EgoVQA [25], EgoTaskQA [11], EgoSchema [12], EgoThink [26], EgoTempo [13], EgoCross [14], EgoBlind [6], EgoMemoria [27], HourVideo [28], EgoLifeQA [1] with different focuses. However, nearly all of them are confined to daytime or well-lit scenarios, leaving model performance in low-light or nighttime conditions largely unexplored. The Oxford Day-and-Night [18] is a partial exception but was not designed for VQA and lacks day-night alignment. This makes EgoNight and EgoNight-VQA fundamentally distinct from prior benchmarks.

2.2 MLLMs for Video Understanding

The rapid development of multimodal large language models (MLLMs) has substantially advanced the frontier of video understanding [29, 30]. Prominent open-source models include Qwen-VL [31], InternVL [32], Video-LLaMA [33], LLaVA-NeXT-Video [34], and GLM-V [35], while closed-source commercial systems such as GPT-4V [36] and Gemini [37] demonstrate even stronger capabilities in video captioning, summarization, and open-ended visual question answering. Building on these advances, egocentric MLLMs have emerged to adapt foundation models from exocentric to first-person perspectives. Representative examples include EgoVLPv2 [15] for improved video-language cross-modal fusion, EgoGPT [1] fine-tuned with egocentric captioning and QA, MM-Ego [27] with a memory mechanism for long videos, and Exo2Ego [16] leveraging exocentric data for egocentric

generalization. These works highlight the potential of MLLMs as egocentric assistants. However, nearly all of them are developed and tested under well-lit daytime conditions, leaving their robustness in low-light or nighttime scenarios unexplored. Nevertheless, almost all existing MLLMs are developed and evaluated under well-lit daytime conditions, with little consideration of low-light or nighttime videos, leaving their robustness in low-light or nighttime scenarios unexplored.

2.3 Cross-Domain Generalization

Domain generalization [38] is a long-standing challenge in computer vision, where models trained on one distribution must adapt to another. Shifts can arise from semantic drift, style changes, or variations in weather and lighting. Many algorithms have been validated across tasks such as image classification [39–43], object detection [44–47], action recognition [48–50], few-shot learning [51–60], autonomous driving [61–66], scene editing and generation [67–70]. In contrast, cross-domain transfer for MLLMs, especially in video understanding, remains underexplored, with only a few recent attempts (e.g., CL-CrossVQA [71], VQA-GEN [72], Super-CLEVR [73]). However, none of them are targeted for egocentric video, which is naturally different from exocentric videos in terms of recorded images, camera motion, and contained information. The most relevant effort to us is EgoCross [14], an egocentric VQA benchmark that moves beyond daily activities to evaluate model generalization across distinct domains such as surgery and industrial settings. In this paper, however, we investigate MLLMs from a different perspective, robustness under nighttime conditions, a critical yet previously overlooked dimension of domain generalization in egocentric video understanding.

3 EgoNight Dataset & Benchmarks

3.1 Video Source Collection

Overview & Design Principles. EgoNight is built to systematically evaluate MLLMs under challenging nighttime conditions, which are critical for developing robust intelligent assistants. The collection of video sources follows four principles: ① *Reflect real-world challenges*, such as walking on dimly lit streets or navigating indoors during power outages; ② *Involve natural camera movements* and preferably capture *actions and interactions* with the environment to evaluate both static perception and dynamic understanding; ③ *Ensure diversity of scenarios, illumination, and task difficulties*, spanning indoor, outdoor, office, and grocery settings, lighting from streetlights, flashlights, headlights, and candles, and task levels from easy (relatively clear), through medium (partially visible), to hard (barely visible). ④ *Enable rigorous analysis through day–night paired videos*, where scenes, trajectories, and actions remain consistent across conditions so that differences can be attributed solely to illumination. To meet these requirements, EgoNight integrates three complementary video sources, as illustrated in the upper part of Fig. 2 and detailed below.

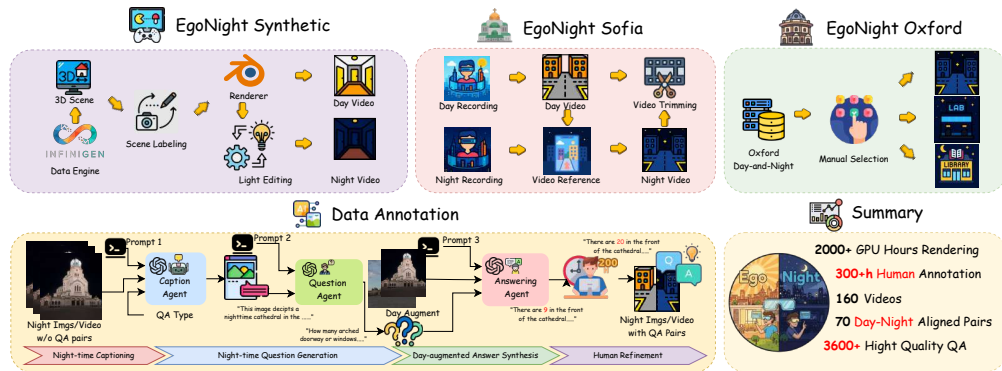


Figure 2: **EgoNight construction and EgoNight-VQA annotation.** EgoNight integrates EgoNight-Synthetic, EgoNight-Sofia, and EgoNight-Oxford sources. Annotation is achieved via a novel three-stage day-augmented Auto QA generation pipeline with 300+ hours of human refinement, resulting in over 3600 high-quality QA pairs.

EgoNight-Synthetic. To achieve perfectly aligned day–night video pairs, we came up with the idea of leveraging the simulation environment. Simulation provides fully controlled environments, allowing us to precisely control scene layout, camera trajectory, and light sources, such that daytime and nighttime counterparts achieve pixel-level and frame-level correspondence, with illumination being the only varying factor. Specifically, we first employ Infinigen [74] as the data engine to generate diverse 3D indoor assets. Human annotators then refine the scenes by correcting errors and removing inappropriate objects before virtually exploring the apartment at an average human walking speed of 1.2 m/s. The resulting camera trajectory is recorded and later replayed for both daytime and nighttime rendering. Finally, we use Blender [17] to render the daytime videos, and adjust the lighting conditions to produce the corresponding nighttime versions.

In total, EgoNight-Synthetic contains 50 pairs of egocentric videos, covering more than 100 environment assets. These include indoor scenes such as kitchens, bathrooms, and living rooms, populated with over 50 diverse object categories (e.g., windows, tables, beds, chairs, lamps, bookshelves, plates). We design multiple illumination setups, ranging from uniformly lit rooms to sparsely localized lighting, and incorporate three difficulty levels with a different range of motion blur, sensor noise, and illumination level. Besides RGB frames, Blender also allows us to generate ground-truth depth and normals (see Appendix Sec. A.1), making EgoNight-Synthetic richer and more versatile.

EgoNight-Sofia. To compensate for the lack of dynamic events and human–environment interactions in synthetic videos, we additionally record our own day–night aligned video pairs, capturing realistic cases that occur in daily life. Although recording strictly aligned pairs is inherently challenging, we design a practical *video-guided recording* strategy combined with post-trimming to achieve decent alignment. Specifically, we first record a daytime video with the ego-wearer exploring an environment while previewing the live feed on a phone screen. For the nighttime counterpart, the same ego-wearer, device setup, and camera perspective are kept fixed. Instead of previewing the live feed, the daytime video is replayed on the phone, serving as guidance for reproducing the same walking speed, viewpoints, and actions. This procedure, after training the ego-wearer to practice following the reference video, proves to be more stable and reliable than alternatives such as setting landmarks or memorizing trajectories, yielding reasonably aligned day–night recordings. Further post-trimming is applied to refine spatial and temporal consistency.

In total, EgoNight-Sofia consists of 20 day–night pairs recorded in Sofia, Bulgaria. Although modest in scale, it is a rare and valuable resource that provides real-world day–night aligned videos capturing diverse daily activities. The dataset spans apartments, offices, streets, neighborhoods, tourist sites, grocery shops, and outdoor fitness areas, with real actions such as drinking water, locking doors, putting down keys, charging devices, and checking price labels in shops. These scenarios naturally yield valuable VQA cases (e.g., “Where did I put my keys?”, “How much is the drink I saw?”, “Did I turn left?”). Illumination sources include street lights, flashlights, small lamps, and candles.

EgoNight-Oxford. We note that Oxford Day–Night [18] is one exception that also contains egocentric videos captured under both daytime and nighttime conditions. Originally proposed for benchmarking 3D vision tasks such as novel view synthesis, the dataset was collected in Oxford and centers on five representative locations. Although it includes illumination variations, the videos are not aligned across day and night. Nevertheless, they provide additional realistic nighttime recordings that increase the scale and diversity of EgoNight, particularly for urban outdoor scenes. Thus, we manually select 20 nighttime segments to form EgoNight-Oxford, following two criteria: i) minimal overlap in locations and trajectories, and ii) dim lighting conditions. These segments serve as a supplementary testbed to assess model generalization under illumination changes in cases where paired day–night alignment is unavailable.

Each EgoNight-Sofia and EgoNight-Oxford video is labeled as easy, medium, or hard by human annotators. Together, these sources provide EgoNight with a balanced mix of precise alignment, natural dynamics, and broad coverage.

3.2 EgoNight-VQA Benchmark Reconstruction

QA Task Taxonomy. To thoroughly assess models from multiple perspectives, we define a diverse taxonomy of 12 QA tasks. Some of these categories are well-studied and have been explored in previous egocentric VQA benchmarks, such as object/action/text recognition, counting, and spatial reasoning. Others are much less studied or newly proposed in EgoNight-VQA, including scene

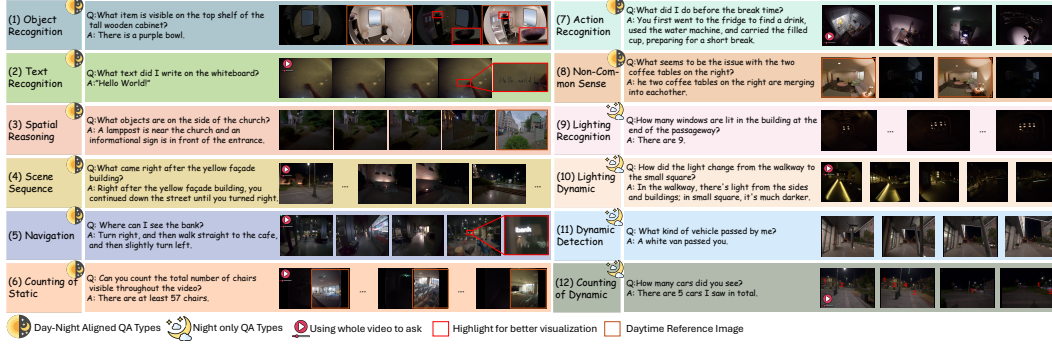


Figure 3: **QA types with examples.** The first eight are *paired* types, where the same question–answer applies to both day and night clips; the last four are *unpaired*, evaluated only at night. QA Types have various durations, with static or spatial tasks (e.g., 1 and 3) using short clips, while dynamic or temporal tasks (e.g., 4 and 5) use full videos.

sequence and navigation (which require not only visual perception but also memory and spatial reasoning), illumination recognition and illumination change (designed to test models’ understanding of lighting concepts), and non–common–sense reasoning (e.g., detecting abnormal cases such as a door inserted into a wall in the synthetic data). More detailed explanations of QA types can be found in the Appendix Sec. A.4.1. We further organize these categories into *paired* and *unpaired* QA types, depending on whether the same questions can be consistently applied across day–night counterparts:

- **Paired QA Types.** These cover contexts that remain unchanged across day and night, allowing the same QA pairs to be used for both videos and thus providing a clean testbed for measuring performance gaps. Specifically, we include: ① *object recognition*, ② *text recognition*, ③ *spatial reasoning*, ④ *scene sequence*, ⑤ *navigation*, ⑥ *counting of static*, ⑦ *action recognition*, and ⑧ *non–common–sense reasoning*.
- **Unpaired QA Types.** These include categories that are impractical to pair across day and night, or are only meaningful in the nighttime condition. We consider: ① *lighting recognition*, ② *lighting dynamic*, ③ *dynamic detection*, and ④ *counting of dynamic*. We control QA clip duration by task type. For static or spatial tasks (e.g., object recognition, lighting recognition), we use short clips of 3 seconds to minimize redundancy; For dynamic or temporal tasks (e.g., action recognition, navigation), the entire video is used to capture the complete context. Following recent works [13, 6], we adopt the *open-ended QA* setting over the closed-form multiple-choice format, as it better reflects natural human–AI interactions.

A detailed summary of each QA type, including whether it is paired or unpaired, clip duration, and example questions, is provided in Fig. 3. This taxonomy makes EgoNight-VQA not only diverse and well-structured but also novel, introducing illumination reasoning and other challenges uniquely tied to nighttime egocentric vision.

Day-Augmented Auto QA Generation. Constructing large-scale QA pairs for nighttime videos is particularly challenging due to low visibility, which makes direct annotation both time-consuming and error-prone. To address this, as illustrated in the lower part of Fig. 2, we design a novel three-stage *day-augmented auto QA generation* pipeline that leverages aligned daytime videos as a strong prior for annotating their nighttime counterparts.

Specifically, the pipeline is tailored to each QA type and consists of three stages:

1) **Nighttime captioning.** For each clip, we prompt advanced MLLMs to generate detailed captions with an explicit focus on the target QA type (e.g., highlighting object-related attributes for object recognition or text/logos for text recognition). This ensures that the captions capture the most key information or construct relevant QA pairs.

2) **Nighttime question generation.** The caption, together with the corresponding night clip, is then fed into the same MLLM to produce diverse question candidates centered on the given QA type. This step encourages variety in phrasing and perspective while maintaining fidelity to the visual content.

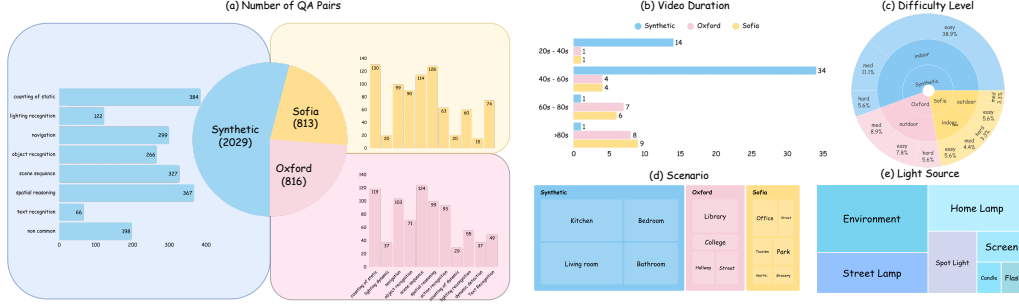


Figure 4: **Statistics of EgoNight-VQA benchmark.** (a) Distribution of QA pairs across QA types and sources. (b) Video duration distribution. (c) Task difficulty levels cross scenarios. (d) Scenario coverage. (e) Illumination coverage.

3) **Day-augmented pseudo answer synthesis.** For paired QA types, pseudo answers are generated by consulting the aligned daytime clip, where content is more visible and less ambiguous. For unpaired QA types or datasets without alignment (e.g., EgoNight-Oxford), answers are instead derived directly from the nighttime clip.

All three stages are powered by GPT-4.1. Empirically, we find that both the QA-type-specific prompting and the inclusion of daytime videos substantially improve the quality and reliability of the generated QA pairs.

Human Annotator Refinement. Finally, human annotators refine QA pairs via three operations: i) **delete**, when QA pairs are meaningless, vague, duplicated, or inconsistent across day-night counterparts (for paired QA types); ii) **modify**, when the question is valid but the answer is wrong (or vice versa), or to resolve ambiguity; iii) **add**, when many pairs are removed or when important, challenging questions, especially about dynamic concepts, are missing.

After the first labeling round, we performed a random double-check to refine low-quality annotations. Thus, although our pipeline combines model generation with human refinement, *every QA pair (3,658 in total) is manually verified at least once*. In total, ~ 200 hours of human effort were invested, ensuring the quality and reliability of EgoNight-VQA.

Dataset Statistics. EgoNight-VQA comprises 3,658 high-quality, fully human-verified QA pairs across 12 task types, sourced from EgoNight-Synthetic, EgoNight-Sofia, and EgoNight-Oxford, with an average of 40 pairs per video. Detailed statistics on QA distribution, video durations, task difficulties, scenarios, and illumination are shown in Fig. 4. Overall, EgoNight-VQA provides a diverse and comprehensive benchmark for evaluating egocentric vision models under the challenging nighttime conditions.

3.3 Benchmarks Beyond Egocentric VQA

Day-Night Correspondence Retrieval. To further assess model capabilities beyond VQA, we introduce *day-night correspondence retrieval*, which evaluates a model’s ability to match visual content across illumination conditions. Specifically, we define two subtasks: i) **Spatial Retrieval (Scene Recognition).** Spatial retrieval, or scene recognition, is a long-standing vision task [75–77]. Here, it is extended: given a query clip and a set of N candidate clips of equal duration s , the model must retrieve the one depicting the same scene. This evaluates a model’s ability to capture and relate spatial relations in egocentric videos, e.g., distinguishing a bedroom from a bathroom or another bedroom. We built this benchmark with 1000 randomly generated meta-tasks. Each task samples a query clip, and the candidate set includes its temporally aligned counterpart (with a temporal shift for added difficulty) plus $N - 1$ negatives from other scenes. Performance is measured by Top-1 accuracy across all tasks. In our setup, we use $N = 10$, $s = 10$ seconds, and a temporal shift of [10, 20] frames. Both *Day (query) $\rightarrow$ Day (database)* and *Day $\rightarrow$ Night* settings are evaluated.

ii) **Temporal Localization.** We further design a temporal localization task to test whether models can align video segments across dynamics. Given a query clip of duration s , the model must localize it within the corresponding full video by predicting its start and end timestamps (t_i, t_j), directly evaluating temporal reasoning (e.g., grounding “The door is being locked” to 10–20s). We construct

1000 meta-tasks, each generated by randomly sampling one clip from its parent full video that is also randomly selected. Inspired by temporal grounding literature [78], we adopt mean Intersection-over-Union (mIoU) between the predicted interval (t_i, t_j) and the ground-truth interval (t_i^*, t_j^*) as the evaluation metric. Consistent with spatial retrieval, we set $s = 10$ seconds and evaluate both $Day \rightarrow Day$ and $Night \rightarrow Day$ settings.

Egocentric Depth Estimation at Night. Depth estimation is a fundamental component of computer vision. On the one hand, extensive research [79–81] has focused on depth estimation in non-egocentric settings (typically not with fisheye cameras), while egocentric depth estimation remains largely underexplored, especially under nighttime conditions. On the other hand, recent works [82, 83] suggest that incorporating depth can enhance models’ spatial reasoning abilities. These two observations motivate us to construct an auxiliary benchmark for *egocentric depth estimation at night*. Specifically, we use EgoNight-Synthetic as the testbed, where ground-truth depth maps are provided by the rendering engine. Thanks to the day–night aligned design, we can quantitatively evaluate models under both controlled daytime and nighttime conditions. For evaluation, we adopt standard depth estimation metrics, including absolute relative error (AbsRel), $\delta_1(1.25)$, $\delta_2(1.25^2)$, and $\delta_3(1.25^3)$, where δ_k measures the percentage of predicted pixels whose relative error is within a threshold of 1.25^k .

4 Experiments

4.1 Evaluated MLLMs & Metrics

We evaluate a broad set of state-of-the-art MLLMs on the proposed benchmarks. i) For **EgoNight-VQA**, we include two closed-source commercial models, GPT-4.1 [36] and Gemini 2.5 Pro [37]; eight open-source models, Qwen2.5-VL (3B, 7B, 72B) [31], VideoLLaMA3 (7B) [33], InternVL3 (8B) [32], GLM-4.1V (9B-Base) [35], and LLaVA-NeXT-Video (7B) [34]; as well as EgoGPT [1], one of the few open-source egocentric models tailored for open-ended QA. Following prior work [13, 25], we adopt an *LLM-as-a-Judge* strategy to assess semantic consistency between predictions and ground truth, and report average accuracy across the test sets. ii) For **day–night correspondence retrieval**, we benchmark feature-based retrieval methods, DINOv2 [84] and Perception Encoder (Percep. Enc.) [85], alongside MLLM-based methods, GPT-4.1 and InternVL3 (8B). As described in Sec. 3.3, Top-1 accuracy (Acc-R@1, %) and mIoU (%) are used for evaluating the spatial and temporal subtasks, respectively. iii) For **egocentric depth estimation**, we test a general monocular depth model (Depth Anything [79, 80]), a 3D reconstruction-based method (VGGTStream [86, 81]), and two egocentric fisheye-specific models (DAC [87] and UniK3D [88]). For Depth Anything and VGGTStream, input fisheye RGB frames and depth maps are undistorted prior to inference for fair comparison. Additional implementation details (e.g., fps for frame extraction, prompts, and model settings) are provided in the Appendix Sec. A.5.

4.2 Results on EgoNight-VQA

The main results of all MLLMs are shown in Tab. 1. In addition, we provide per-QA performance comparisons between day (striped bars) and night (solid bars) for paired QA types (Fig. 5(a)) and report nighttime performance across all QA types (Fig. 5(b)), based on averages across all models. Note that non-common case detection is available only in EgoNight-Synthetic, while dynamic events and actions are included only in the real-world data.

From the results in Tab. 1, we observe that almost all MLLMs struggle on our benchmark, with maximum averaged accuracies of 30.93% from the closed-source GPT-4.1, 20.06% from the open-source InternVL3-8B, and 14.29% from the egocentric EgoGPT. Fig. 5(a) further highlights the performance gap, showing declines of 32.8% and 25.0% on EgoNight-Synthetic and EgoNight-Sofia, respectively. Together, these results underscore the substantial challenges posed by our benchmark, exposing the limitations of current MLLMs under nighttime scenarios and highlighting the need for more illumination-robust models. Beyond the overall trends, we note three additional insights from Tab. 1: i) Closed-source models perform best. Within open-source models, Qwen2.5-VL generally improves with scale, yet InternVL outperforms the larger Qwen2.5-VL (72B), suggesting that size alone is insufficient. The relatively low results of EgoGPT further emphasize the need for more robust egocentric models. ii) EgoNight-Oxford achieves the highest scores, but its illumination conditions are more challenging than those in EgoNight-Synthetic and EgoNight-Sofia (Sec. A.4.2, Appendix).

Models	EgoNight-Synthetic			EgoNight-Sofia			EgoNight-Oxford			Avg. -
	Easy	Medium	Hard	Easy	Medium	Hard	Easy	Medium	Hard	
Closed-Source MLLMs										
GPT-4.1	29.30	26.87	18.87	32.04	29.35	31.69	39.72	37.13	40.72	30.93
Gemini 2.5 Pro	31.05	24.81	16.51	38.24	26.81	28.87	36.75	36.81	27.88	30.60
Open-source MLLMs										
InternVL3-8B	20.21	15.50	16.98	24.03	21.74	20.42	22.90	20.85	16.36	20.06
Qwen2.5-VL-72B	18.39	15.25	12.26	24.03	17.03	20.42	24.81	22.80	16.36	18.99
Qwen2.5-VL-7B	13.01	13.95	13.68	15.44	12.68	12.68	13.74	13.36	12.73	13.44
Qwen2.5-VL-3B	14.69	10.34	7.08	15.50	13.04	12.68	17.18	11.40	12.12	13.41
GLM-4.1V-9B-Base	19.09	13.70	15.57	18.60	18.48	16.20	17.15	22.15	18.79	18.20
VideoLLaMA3-7B	16.85	13.44	14.62	11.11	10.87	9.15	12.26	10.46	9.15	13.64
LLaVA-NeXT-Video-7B	6.36	11.37	1.89	13.95	9.78	14.79	3.05	2.61	3.03	7.28
Egocentric MLLMs										
EgoGPT	15.79	13.55	12.04	12.41	12.13	10.36	12.37	13.58	13.68	14.29

Table 1: **Comparison results on EgoNight-VQA.** Accuracies (%) of OpenQA results across three datasets and three difficulty levels. We compare closed-source models, open-source models, and egocentric-specific models.

This indicates that without paired day videos and our day-augmented auto-labeling strategy, even human annotators face difficulties generating challenging QA pairs, underscoring the practical value of our dataset design; iii) Overall, performance declines across task levels (easy, medium, hard), validating the diversity and difficulty of our benchmark.

From the per-QA results in Fig. 5(a) and Fig. 5(b), we further observe three key trends: i) Models perform better on perception-oriented tasks (e.g., object recognition, text recognition, scene sequence) than reasoning-oriented tasks (e.g., navigation, counting, non-common-sense reasoning cases) under daytime conditions. However, at night, perception tasks suffer larger performance drops, indicating their higher sensitivity to illumination, whereas reasoning tasks, though harder overall, are relatively less affected since they rely more on temporal and contextual cues. ii) MLLMs achieve substantially lower accuracy on our newly proposed tasks, such as lighting recognition, lighting dynamics, scene sequence, dynamic detection, navigation, and non-common-sense reasoning, suggesting that existing MLLMs generalize poorly to novel tasks compared with well-studied ones like object recognition. iii) Each dataset in Fig. 5(b) emphasizes distinct aspects of nighttime challenges, together providing complementary perspectives that ensure EgoNight spans a balanced range of perception–reasoning difficulties under low-light conditions.

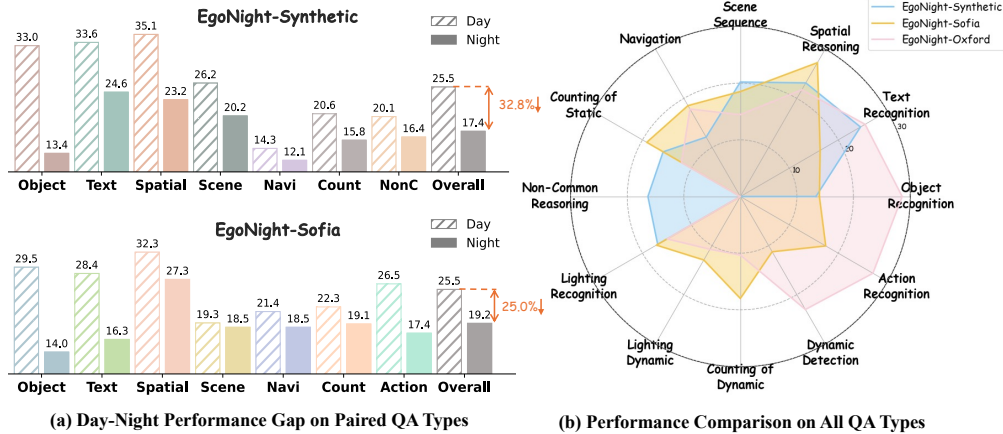


Figure 5: **Performance analysis of MLLMs on EgoNight-VQA.** (a) Day–night performance gap across paired QA types, showing consistent degradation at night. (b) Nighttime performance across all 12 QA types.

4.3 Results on Day-Night Correspondence Retrieval

The results of day–night retrieval are reported in Tab. 2. The gap between Night–Day and Day–Day shows that cross-illumination retrieval remains highly challenging compared with in-domain retrieval. For spatial retrieval, GPT-4.1 consistently outperforms other methods, achieving over 80% accuracy. This suggests that Retrieval-Augmented Generation methods could further improve performance, as Fig. 5(a) already shows that daytime inputs significantly benefit the models. For temporal retrieval, however, GPT-4.1, despite its strong results on egocentric VQA (Tab. 1) and spatial retrieval, shows a substantial drop compared with feature-based methods (DINOv2 and Perception Encoder). A similar degradation is observed for InternVL3-8B. These findings suggest that while MLLMs excel at spatial semantic understanding, they struggle with temporal reasoning, such as timestamp prediction, which is critical for temporal localization. Further results on temporal limitations are provided in Appendix A.6.

Models	Spatial Retrieval (Acc - R@1 % $\uparrow$)				Temporal Localization (mIoU % $\uparrow$)			
	EgoNight-Synthetic		EgoNight-Sofia		EgoNight-Synthetic		EgoNight-Sofia	
	Day→Day	Night→Day	Day→Day	Night→Day	Day→Day	Night→Day	Day→Day	Night→Day
DINOv2	45.7	28.7	84.5	74.5	-	33.7	-	33.1
Percep. Enc.	65.4	41.6	89.8	80.9	-	32.9	-	33.4
GPT-4.1	75.6	54.1	92.5	84.5	14.7	10.0	21.2	15.5
InternVL3-8B	39.4	27.7	73.9	56.3	10.2	9.9	12.5	13.3

Table 2: **Night-to-Day retrieval performance.** Each dataset is evaluated on both Day→Day and Night→Day settings.

4.4 Results on Depth Estimation

Results for depth estimation are reported in Tab. 3. The relatively low scores across all models highlight the difficulty of our EgoNight dataset, which combines egocentric motion, complex geometry, and extreme lighting variations. A clear gap between daytime and nighttime performance again underscores the challenges of low-light conditions. Among the methods, fisheye-based methods (DAC and UniK3D) outperform general depth estimators, suggesting the need for egocentric-specific algorithms. Additional qualitative results are provided in Sec. A.7.3.

Method	Abs Rel $\downarrow$		δ_1 (1.25) $\uparrow$		δ_2 (1.25 ²) $\uparrow$		δ_3 (1.25 ³) $\uparrow$	
	Day	Night	Day	Night	Day	Night	Day	Night
Depth Anything (U)	0.297	0.302	0.249	0.237	0.463	0.447	0.622	0.60
VGGTStream (U)	0.293	0.298	0.234	0.232	0.447	0.442	0.615	0.609
DAC (F)	0.245	0.292	0.255	0.216	0.495	0.425	0.684	0.602
UniK3D (F)	0.224	0.253	0.280	0.254	0.524	0.481	0.706	0.658

Table 3: Depth estimation results on EgoNight-Synthetic. **U**: undistorted input; **F**: fisheye input.

5 Conclusion

In this work, we introduced EgoNight, the first benchmark suite designed to systematically evaluate egocentric multimodal large language models (MLLMs) under challenging nighttime conditions. EgoNight integrates synthetic and real-world videos with day–night alignment, enabling rigorous analysis of illumination effects. Building upon this data, we proposed EgoNight-VQA, spanning 12 QA types with 3,658 human-verified pairs, alongside two complementary benchmarks: day–night correspondence retrieval and egocentric depth estimation. Experiments reveal that even state-of-the-art MLLMs struggle under low-light conditions, with performance dropping substantially compared to daytime. This highlights that nighttime egocentric vision remains far from being solved, motivating future research into illumination-robust egocentric perception and reasoning. We believe EgoNight provides a valuable and timely benchmark that will drive progress toward more reliable egocentric AI assistants.

References

- [1] J. Yang, S. Liu, H. Guo, Y. Dong, X. Zhang, S. Zhang, P. Wang, Z. Zhou, B. Xie, Z. Wang, *et al.*, “Egolife: Towards egocentric life assistant,” in *CVPR*, 2025.
- [2] Y. Fu, C. Wang, Y. Fu, Y.-X. Wang, C. Bai, X. Xue, and Y.-G. Jiang, “Embodied one-shot video recognition: Learning from actions of a virtual embodied agent,” in *ACM Multimedia*, 2019.
- [3] K. Li, Q. Xu, T. Qian, Y. Fu, Y. Jiao, and X. Wang, “Clivis: Unleashing cognitive map through linguistic-visual synergy for embodied visual reasoning,” *arXiv preprint arXiv:2506.17629*, 2025.
- [4] S. Kareer, D. Patel, R. Punamiya, P. Mathur, S. Cheng, C. Wang, J. Hoffman, and D. Xu, “Egomimic: Scaling imitation learning via egocentric video,” in *ICRA*, 2025.
- [5] T. Qian, J. Chen, L. Zhuo, Y. Jiao, and Y.-G. Jiang, “Nuscenes-qa: A multi-modal visual question answering benchmark for autonomous driving scenario,” in *AAAI*, 2024.
- [6] J. Xiao, N. Huang, H. Qiu, Z. Tao, X. Yang, R. Hong, M. Wang, and A. Yao, “Egoblind: Towards egocentric visual assistance for the blind people,” *arXiv preprint arXiv:2503.08221*, 2025.
- [7] G. Zhang, D. Zhang, L. Duan, and G. Han, “Accessible robot control in mixed reality,” 2023.
- [8] D. Damen, H. Doughty, G. M. Farinella, S. Fidler, A. Furnari, E. Kazakos, D. Moltisanti, J. Munro, T. Perrett, W. Price, *et al.*, “The epic-kitchens dataset: Collection, challenges and baselines,” *TPAMI*, 2020.
- [9] K. Grauman, A. Westbury, E. Byrne, Z. Chavis, A. Furnari, R. Girdhar, J. Hamburger, H. Jiang, M. Liu, X. Liu, *et al.*, “Ego4d: Around the world in 3,000 hours of egocentric video,” in *CVPR*, 2022.
- [10] K. Grauman, A. Westbury, L. Torresani, K. Kitani, J. Malik, T. Afouras, K. Ashutosh, V. Baiyya, S. Bansal, B. Boote, *et al.*, “Ego-exo4d: Understanding skilled human activity from first-and third-person perspectives,” in *CVPR*, 2024.
- [11] B. Jia, T. Lei, S.-C. Zhu, and S. Huang, “Egotaskqa: Understanding human tasks in egocentric videos,” *NeurIPS*, 2022.
- [12] K. Mangalam, R. Akshulakov, and J. Malik, “Egoschema: A diagnostic benchmark for very long-form video language understanding,” *NeurIPS*, 2023.
- [13] C. Plizzari, A. Tonioni, Y. Xian, A. Kulshrestha, and F. Tombari, “Omnia de egotempo: Benchmarking temporal understanding of multi-modal llms in egocentric videos,” in *CVPR*, 2025.
- [14] Y. Li, Y. Fu, T. Qian, Q. Xu, S. Dai, D. P. Paudel, L. Van Gool, and X. Wang, “Egocross: Benchmarking multimodal large language models for cross-domain egocentric video question answering,” *arXiv preprint arXiv:2508.10729*, 2025.
- [15] S. Pramanick, Y. Song, S. Nag, K. Q. Lin, H. Shah, M. Z. Shou, R. Chellappa, and P. Zhang, “Egovlpv2: Egocentric video-language pre-training with fusion in the backbone,” in *ICCV*, 2023.
- [16] H. Zhang, Q. Chu, M. Liu, Y. Wang, B. Wen, F. Yang, T. Gao, D. Zhang, Y. Wang, and L. Nie, “Exo2ego: Exocentric knowledge guided mllm for egocentric video understanding,” *arXiv preprint arXiv:2503.09143*, 2025.
- [17] B. Iraci, *Blender cycles: lighting and rendering cookbook*. Packt Publishing Ltd, 2013.
- [18] Z. Wang, W. Bian, X. Li, Y. Tao, J. Wang, M. Fallon, and V. A. Prisacariu, “Seeing in the dark: Benchmarking egocentric 3d vision with the oxford day-and-night dataset,” *arXiv preprint arXiv:2506.04224*, 2025.
- [19] Y. Huang, G. Chen, J. Xu, M. Zhang, L. Yang, B. Pei, H. Zhang, L. Dong, Y. Wang, L. Wang, *et al.*, “Egoexolearn: A dataset for bridging asynchronous ego-and exo-centric view of procedural activities in real world,” in *CVPR*, 2024.

- [20] S. Sudhakaran, S. Escalera, and O. Lanz, “Lsta: Long short-term attention for egocentric action recognition,” in *CVPR*, 2019.
- [21] X. Ren and C. Gu, “Figure-ground segmentation improves handled object recognition in egocentric video,” in *CVPR*, 2010.
- [22] Z. Luo, R. Hachiuma, Y. Yuan, and K. Kitani, “Dynamics-regulated kinematic policy for egocentric pose estimation,” *NeurIPS*, 2021.
- [23] G. Liu, H. Tang, H. M. Latapie, J. J. Corso, and Y. Yan, “Cross-view exocentric to egocentric video synthesis,” in *ACM Multimedia*, 2021.
- [24] Y. Fu, R. Wang, Y. Fu, D. P. Paudel, X. Huang, and L. Van Gool, “Objectrelator: Enabling cross-view object relation understanding in ego-centric and exo-centric videos,” *ICCV*, 2025.
- [25] C. Fan, “Egovqa-an egocentric video question answering benchmark dataset,” in *ICCV Workshop*, 2019.
- [26] S. Cheng, Z. Guo, J. Wu, K. Fang, P. Li, H. Liu, and Y. Liu, “Egothink: Evaluating first-person perspective thinking capability of vision-language models,” in *CVPR*, 2024.
- [27] H. Ye, H. Zhang, E. Daxberger, L. Chen, Z. Lin, Y. Li, B. Zhang, H. You, D. Xu, Z. Gan, *et al.*, “Mm-ego: Towards building egocentric multimodal llms for video qa,” *arXiv preprint arXiv:2410.07177*, 2024.
- [28] K. Chandrasegaran, A. Gupta, L. M. Hadzic, T. Kota, J. He, C. Eyzaguirre, Z. Durante, M. Li, J. Wu, and L. Fei-Fei, “Hourvideo: 1-hour video-language understanding,” *NeurIPS*, 2024.
- [29] D. Wen, K. Peng, J. Zheng, Y. Chen, Y. Shi, J. Wei, R. Liu, K. Yang, and R. Stiefelhofen, “Mica: Multi-agent industrial coordination assistant,” *arXiv preprint arXiv:2509.15237*, 2025.
- [30] Y. Lyu, X. Zheng, J. Zhou, and L. Wang, “Unibind: Llm-augmented unified and balanced representation space to bind them all,” in *CVPR*, 2024.
- [31] J. Bai, S. Bai, S. Yang, S. Wang, S. Tan, P. Wang, J. Lin, C. Zhou, and J. Zhou, “Qwen-vl: A versatile vision-language model for understanding, localization, text reading, and beyond,” *arXiv preprint arXiv:2308.12966*, 2023.
- [32] Z. Chen, J. Wu, W. Wang, W. Su, G. Chen, S. Xing, M. Zhong, Q. Zhang, X. Zhu, L. Lu, *et al.*, “Internvl: Scaling up vision foundation models and aligning for generic visual-linguistic tasks,” in *CVPR*, 2024.
- [33] H. Zhang, X. Li, and L. Bing, “Video-llama: An instruction-tuned audio-visual language model for video understanding,” *arXiv preprint arXiv:2306.02858*, 2023.
- [34] B. Li, Y. Zhang, D. Guo, R. Zhang, F. Li, H. Zhang, K. Zhang, P. Zhang, Y. Li, Z. Liu, *et al.*, “Llava-onevision: Easy visual task transfer,” *arXiv preprint arXiv:2408.03326*, 2024.
- [35] W. Hong, W. Yu, X. Gu, G. Wang, G. Gan, H. Tang, J. Cheng, J. Qi, J. Ji, L. Pan, *et al.*, “Glm-4.1 v-thinking: Towards versatile multimodal reasoning with scalable reinforcement learning,” *arXiv e-prints*, 2025.
- [36] J. Achiam, S. Adler, S. Agarwal, L. Ahmad, I. Akkaya, F. L. Aleman, D. Almeida, J. Altenschmidt, S. Altman, S. Anadkat, *et al.*, “Gpt-4 technical report,” *arXiv preprint arXiv:2303.08774*, 2023.
- [37] G. Comanici, E. Bieber, M. Schaekermann, I. Pasupat, N. Sachdeva, I. Dhillon, M. Blistein, O. Ram, D. Zhang, E. Rosen, *et al.*, “Gemini 2.5: Pushing the frontier with advanced reasoning, multimodality, long context, and next generation agentic capabilities,” *arXiv preprint arXiv:2507.06261*, 2025.
- [38] K. Zhou, Z. Liu, Y. Qiao, T. Xiang, and C. C. Loy, “Domain generalization: A survey,” *TPAMI*, 2022.

- [39] D. Li, Y. Yang, Y.-Z. Song, and T. M. Hospedales, “Deeper, broader and artier domain generalization,” in *ICCV*, 2017.
- [40] K. Zhou, Y. Yang, Y. Qiao, and T. Xiang, “Domain generalization with mixstyle,” *arXiv preprint arXiv:2104.02008*, 2021.
- [41] J. Zhang, L. Gao, B. Hao, H. Huang, J. Song, and H. Shen, “From global to local: Multi-scale out-of-distribution detection,” *IEEE TIP*, 2023.
- [42] K. Peng, D. Wen, K. Yang, A. Luo, Y. Chen, J. Fu, M. S. Sarfraz, A. Roitberg, and R. Stiefelhagen, “Advancing open-set domain generalization using evidential bi-level hardest domain scheduler,” *NeurIPS*, 2024.
- [43] K. Peng, D. Wen, S. M. Saquib, Y. Chen, J. Zheng, D. Schneider, K. Yang, J. Wu, A. Roitberg, and R. Stiefelhagen, “Mitigating label noise using prompt-based hyperbolic meta-learning in open-set domain generalization,” *arXiv preprint arXiv:2412.18342*, 2024.
- [44] Y. Fu, Y. Wang, Y. Pan, L. Huai, X. Qiu, Z. Shangguan, T. Liu, Y. Fu, L. Van Gool, and X. Jiang, “Cross-domain few-shot object detection via enhanced open-set object detector,” in *ECCV*, 2024.
- [45] Y. Li, X. Qiu, Y. Fu, J. Chen, T. Qian, X. Zheng, D. P. Paudel, Y. Fu, X. Huang, L. Van Gool, *et al.*, “Domain-rag: Retrieval-guided compositional image generation for cross-domain few-shot object detection,” *NeurIPS*, 2025.
- [46] J. Pan, Y. Liu, X. He, L. Peng, J. Li, Y. Sun, and X. Huang, “Enhance then search: An augmentation-search strategy with foundation models for cross-domain few-shot object detection,” in *CVPR workshop*, 2025.
- [47] Y. Zheng, X. Chen, Y. Zheng, S. Gu, R. Yang, B. Jin, P. Li, C. Zhong, Z. Wang, L. Liu, *et al.*, “Gaussiangrasper: 3d language gaussian splatting for open-vocabulary robotic grasping,” *arXiv preprint arXiv:2403.09637*, 2024.
- [48] B. Pan, Z. Cao, E. Adeli, and J. C. Niebles, “Adversarial cross-domain action recognition with co-attention,” in *AAAI*, 2020.
- [49] W. Bian, D. Tao, and Y. Rui, “Cross-domain human action recognition,” *IEEE Transactions on Systems, Man, and Cybernetics, Part B (Cybernetics)*, 2011.
- [50] K. Peng, J. Fu, K. Yang, D. Wen, Y. Chen, R. Liu, J. Zheng, J. Zhang, M. S. Sarfraz, R. Stiefelhagen, *et al.*, “Referring atomic video action recognition,” in *ECCV*, 2024.
- [51] Y. Fu, Y. Fu, and Y.-G. Jiang, “Meta-fdmixup: Cross-domain few-shot learning guided by labeled target data,” in *ACM Multimedia*, 2021.
- [52] Y. Guo, N. C. Codella, L. Karlinsky, J. V. Codella, J. R. Smith, K. Saenko, T. Rosing, and R. Feris, “A broader study of cross-domain few-shot learning,” in *ECCV*, 2020.
- [53] Y. Fu, Y. Xie, Y. Fu, J. Chen, and Y.-G. Jiang, “Wave-san: Wavelet based style augmentation network for cross-domain few-shot learning,” *arXiv preprint*, 2022.
- [54] J. Zhang, J. Song, L. Gao, and H. Shen, “Free-lunch for cross-domain few-shot learning: Style-aware episodic training with robust contrastive learning,” in *ACM multimedia*, 2022.
- [55] L. Zhuo, Y. Fu, J. Chen, Y. Cao, and Y.-G. Jiang, “Tgdm: Target guided dynamic mixup for cross-domain few-shot learning,” in *ACM Multimedia*, 2022.
- [56] H. Tang, C. Yuan, Z. Li, and J. Tang, “Learning attention-guided pyramidal features for few-shot fine-grained recognition,” *Pattern Recognition*, 2022.
- [57] Y. Fu, Y. Xie, Y. Fu, J. Chen, and Y.-G. Jiang, “Me-d2n: Multi-expert domain decompositional network for cross-domain few-shot learning,” in *ACM Multimedia*, 2022.
- [58] H. Tang, J. Liu, S. Yan, R. Yan, Z. Li, and J. Tang, “M3net: Multi-view encoding, matching, and fusion for few-shot fine-grained action recognition,” in *ACM Multimedia*, 2023.

- [59] L. Zhuo, Z. Wang, Y. Fu, and T. Qian, “Prompt as free lunch: Enhancing diversity in source-free cross-domain few-shot learning through semantic-guided prompting,” *arXiv preprint*, 2024.
- [60] Y. Fu, Y. Xie, Y. Fu, and Y.-G. Jiang, “Styleadv: Meta style adversarial training for cross-domain few-shot learning,” in *CVPR*, 2023.
- [61] G. Li, Z. Ji, X. Qu, R. Zhou, and D. Cao, “Cross-domain object detection for autonomous driving: A stepwise domain adaptative yolo approach,” *T-IV*, 2022.
- [62] J. Li, R. Xu, J. Ma, Q. Zou, J. Ma, and H. Yu, “Domain adaptive object detection for autonomous driving under foggy weather,” in *WACV*, 2023.
- [63] Z. Wu, T. Liu, L. Luo, Z. Zhong, J. Chen, H. Xiao, C. Hou, H. Lou, Y. Chen, R. Yang, Y. Huang, X. Ye, Z. Yan, Y. Shi, Y. Liao, and H. Zhao, “Mars: An instance-aware, modular and realistic simulator for autonomous driving,” *CICAI*, 2023.
- [64] X. Zheng, J. Zhu, Y. Liu, Z. Cao, C. Fu, and L. Wang, “Both style and distortion matter: Dual-path unsupervised domain adaptation for panoramic semantic segmentation,” in *CVPR*, 2023.
- [65] X. Zheng, T. Pan, Y. Luo, and L. Wang, “Look at the neighbor: Distortion-aware unsupervised domain adaptation for panoramic semantic segmentation,” in *ICCV*, 2023.
- [66] T. Brödermann, C. Sakaridis, Y. Fu, and L. Van Gool, “Cafuser: Condition-aware multimodal fusion for robust semantic perception of driving scenes,” *IEEE RA-L*, 2025.
- [67] D. Zhang, C. Fernandez-Labrador, and C. Schroers, “Coarf: Controllable 3d artistic style transfer for radiance fields,” in *3DV*, 2024.
- [68] D. Zhang, J. Wang, S. Wang, M. Mihajlovic, S. Prokudin, H. P. Lensch, and S. Tang, “Rise-sdf: A relightable information-shared signed distance field for glossy object inverse rendering,” in *3DV*, 2025.
- [69] H. Yu, D. Zhang, P. Xie, and T. Zhang, “Point-based radiance fields for controllable human motion synthesis,” 2023.
- [70] N. Savov, N. Kazemi, D. Zhang, D. P. Paudel, X. Wang, and L. V. Gool, “Statespacediffuser: Bringing long context to diffusion world models,” 2025.
- [71] Y. Zhang, H. Chen, A. Frikha, D. Krompass, G. Zhang, J. Gu, and V. Tresp, “Cl-cross vqa: A continual learning benchmark for cross-domain visual question answering,” in *WACV*, 2025.
- [72] S. J. Unni, R. Moraffah, and H. Liu, “Vqa-gen: A visual question answering benchmark for domain generalization,” *arXiv preprint arXiv:2311.00807*, 2023.
- [73] Z. Li, X. Wang, E. Stengel-Eskin, A. Kortylewski, W. Ma, B. Van Durme, and A. L. Yuille, “Super-clevr: A virtual benchmark to diagnose domain robustness in visual reasoning,” in *CVPR*, 2023.
- [74] A. Raistrick, L. Lipson, Z. Ma, L. Mei, M. Wang, Y. Zuo, K. Kayan, H. Wen, B. Han, Y. Wang, A. Newell, H. Law, A. Goyal, K. Yang, and J. Deng, “Infinite photorealistic worlds using procedural generation,” in *CVPR*, 2023.
- [75] R. Arandjelovic, P. Gronat, A. Torii, T. Pajdla, and J. Sivic, “Netvlad: Cnn architecture for weakly supervised place recognition,” in *CVPR*, 2016.
- [76] Y. Miao, F. Engelmann, O. Vysotska, F. Tombari, M. Pollefeys, and D. B. Baráth, “Scene-graphloc: Cross-modal coarse visual localization on 3d scene graphs,” in *ECCV*, 2024.
- [77] Y. Shi, R. Yang, Z. Wu, P. Li, C. Liu, H. Zhao, and G. Zhou, “City-scale continual neural semantic mapping with three-layer sampling and panoptic representation,” *Knowledge-Based Systems*, 2024.
- [78] X. Gu, H. Fan, Y. Huang, T. Luo, and L. Zhang, “Context-guided spatio-temporal video grounding,” in *CVPR*, 2024.

- [79] L. Yang, B. Kang, Z. Huang, X. Xu, J. Feng, and H. Zhao, “Depth anything: Unleashing the power of large-scale unlabeled data,” in *CVPR*, 2024.
- [80] L. Yang, B. Kang, Z. Huang, Z. Zhao, X. Xu, J. Feng, and H. Zhao, “Depth anything v2,” *arXiv:2406.09414*, 2024.
- [81] J. Wang, M. Chen, N. Karaev, A. Vedaldi, C. Rupprecht, and D. Novotny, “Vggt: Visual geometry grounded transformer,” in *CVPR*, 2025.
- [82] B. Chen, Z. Xu, S. Kirmani, B. Ichter, D. Sadigh, L. Guibas, and F. Xia, “Spatialvlm: Endowing vision-language models with spatial reasoning capabilities,” in *CVPR*, 2024.
- [83] Y. Liu, M. Ma, X. Yu, P. Ding, H. Zhao, M. Sun, S. Huang, and D. Wang, “Ssr: Enhancing depth perception in vision-language models via rationale-guided spatial reasoning,” *arXiv preprint arXiv:2505.12448*, 2025.
- [84] M. Oquab, T. Darcet, T. Moutakanni, H. V. Vo, M. Szafraniec, V. Khalidov, P. Fernandez, D. Haziza, F. Massa, A. El-Nouby, R. Howes, P.-Y. Huang, H. Xu, V. Sharma, S.-W. Li, W. Galuba, M. Rabbat, M. Assran, N. Ballas, G. Synnaeve, I. Misra, H. Jegou, J. Mairal, P. Labatut, A. Joulin, and P. Bojanowski, “Dinov2: Learning robust visual features without supervision,” 2023.
- [85] D. Bolya, P.-Y. Huang, P. Sun, J. H. Cho, A. Madotto, C. Wei, T. Ma, J. Zhi, J. Rajasegaran, H. Rasheed, J. Wang, M. Monteiro, H. Xu, S. Dong, N. Ravi, D. Li, P. Dollár, and C. Feichtenhofer, “Perception encoder: The best visual embeddings are not at the output of the network,” 2025.
- [86] D. Zhuo, W. Zheng, J. Guo, Y. Wu, J. Zhou, and J. Lu, “Streaming 4d visual geometry transformer,” *arXiv preprint arXiv:2507.11539*, 2025.
- [87] Y. Guo, S. Garg, S. M. H. Miangoleh, X. Huang, and L. Ren, “Depth any camera: Zero-shot metric depth estimation from any camera,” in *CVPR*, 2025.
- [88] L. Piccinelli, C. Sakaridis, M. Segu, Y.-H. Yang, S. Li, W. Abbeloos, and L. Van Gool, “UniK3D: Universal camera monocular 3d estimation,” in *CVPR*, 2025.
- [89] Y. Zheng, R. Zhang, J. Zhang, Y. Ye, Z. Luo, Z. Feng, and Y. Ma, “Llamafactory: Unified efficient fine-tuning of 100+ language models,” in *ACL*, 2024.

A Appendix

A.1 More Video Source Construction Details

EgoNight-Synthetic Construction. For EgoNight-Synthetic Construction, we first use the coarse progressive generation method with a fast solver in infinigen [74] to generate 3D scenes in Blender format. Then, a human annotator will edit the scene in the following sequence:

- Explore and edit the scene to remove unreasonable cases and make the indoor scene as natural as possible.
- Add light source in the scene if the generated scene does not include enough illumination to create enough illumination gap between the day and night.
- Record camera trajectory by exploring the whole indoor scene.
- Change the camera model and resolution. Set rendering samples and frames. For all synthetic dataset, we use the Blender build-in Panoramic Fisheye Equisolid camera with Lens 10.5 and field of view 180°.
- Create night scene by modifying the light source, motion blur, and environment map.
- Render the day and night pair using Blender [17].

During the dataset construction, we apply home light source during night for 30 scenes and spot light source for 20 scenes to simulate torch light in real life. To create different difficulty levels, we apply different rendering sample size (higher sample size gives lower noise in the final image), spot light size, and motion blur to part of the data as shown in Tab. 4. We also show different modality and difficulty level in Fig. 6

difficulty level	sample size	light condition	motion blur shutter
Easy	4096	105°/ few light on	-
Medium	512	40°-50°spot light / all light off	-
Hard	512	40°-50°spot light / all light off	1-2

Table 4: Difficulty level and corresponding rendering settings.

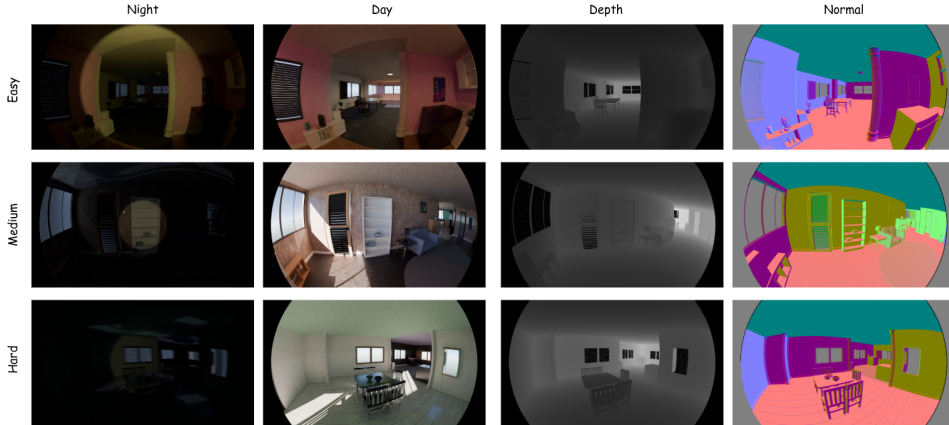


Figure 6: More examples and modalities of synthetic datasets.

EgoNight-Sofia Construction. In total, four participants were involved. The recording setup included three different GoPros, a head-mounted rig to fix the camera on the forehead and mimic human-eye perspective, several phones for live preview or daytime video guidance, and diverse lighting sources such as flashlights, spotlights, and candles. The process followed a video-guided recording strategy, as introduced in Sec. 3.1: the ego-wearer first recorded a daytime video while previewing the live feed on a phone, and for the nighttime counterpart replayed the daytime video on the phone as guidance to replicate the same setup, walking speed, viewpoints, and actions. Videos were collected across a wide range of environments, including indoor scenes (apartments, workplaces, grocery shops, building receptions) and outdoor scenarios (fitness areas, tourist landmarks, and street views). Post-trimming was applied to each day–night pair to further ensure alignment. On average, it took around 2-3 hours to produce one paired data.

EgoNight-Oxford Construction. We credit the contribution of Seeing in the Dark dataset [18], which provides multiple sequences of egocentric videos in the night in a various environment. We built our EgoNight-VQA dataset partially upon this work. Firstly, We enumerated all nighttime clips in Oxford Day–Night and performed a two-stage filtering. (1) Screening for uniqueness of place. We cross-checked scene metadata (route notes/time stamps) to avoid repeated paths within the same landmark. (2) Stratified diversity & quality sampling. Remaining clips were scored on a 1–5 rubric along axes designed for egocentric, low-light evaluation: illumination type (ambient only / mixed artificial / high-contrast point sources), illumination hardness (soft vs. specular/point), exposure stability (auto-gain pumping, blown highlights), scene dynamics (pedestrians/vehicles, occlusions), camera motion pattern (walk, run, head turns), and task context (navigation, road-crossing, object interaction, signage reading). The final set comprises 20 sequences that maximize lighting/task diversity under egocentric night settings while avoiding place overlap and task overlap.

In total, collecting the three video sources required over 100 hours of human effort.

A.2 More Benchmark Implementation Details

As in Sec. 3.2, our auto-labeling pipeline is QA-type specific and involves three customized prompts for captioning, question generation, and answer synthesis. The detailed prompts are shown in Fig. 8.

A.2.1 EgoNight-VQA Human Labeling

We hired several participants to review and refine the QA pairs generated by our three-stage day-augmented auto-labeling pipeline, compensating them at a rate of €20 per video. Each participant was provided with a detailed labeling instruction document and an onboarding meeting to ensure the guidelines were clearly conveyed. A simplified version of the labeling tutorial is included as in Fig. 7.

Annotation Tutorial (Simplified)

Read Me First: Please follow the labeling pipeline carefully and complete each step as instructed. On average, annotating one video takes about 2 hours. Easier cases may take less time, but in general, each video should take more than 1.5 hours to ensure high-quality annotations. (The first video may take longer, as you will need to familiarize yourself with the pipeline.) We will randomly check the labeled data, and annotators will be required to refine their work if the quality does not meet expectations.

Step 1: Preparation. You are expected to first download the paired `day.mp4` and `night.mp4` videos (aligned in time, except for unpaired tasks), together with the QA text file (`.txt`), which contains candidate QAs grouped by QA type (e.g., `counting.txt`). Before starting annotation, you should carefully watch both the daytime and nighttime videos to fully understand the scenario and activities.

Step 2: QA Verification and Refinement. For each QA pair, you should apply one of three operations:

- **Delete:** You should remove QAs that are meaningless, vague, irrelevant, duplicated, or inconsistent between day–night pairs (for paired QA types).
- **Modify:** If the question is reasonable but the answer is incorrect (e.g., counting errors, wrong action duration), you should correct the answer. You may also rephrase the question to eliminate ambiguity (e.g., clarifying “left/right” as relative to the ego-wearer).
- **Add:** If too many pairs are deleted, or if you notice interesting and challenging cases missing, you need to add new QAs. This is especially important for low-frequency tasks, e.g., dynamic detection or counting of dynamics.

Step 3: Special Cases.

- For paired QA types (e.g., object recognition, spatial reasoning), you must ensure the same QAs apply to both day and night videos.
- For unpaired QA types (e.g., lighting recognition, dynamic detection), you only need to ensure correctness on the nighttime video.
- For dynamic events, you are expected to specify temporal spans, e.g., *Q: “Around which time does a red car pass by?” A: “At frames 4–6.”*

Step 4: Post-processing. Once QAs were validated, the answer field should be renamed from “`answer`” to “`human_answer`”.

Appendix: Paired & Non-Paired QA Types. The same as described in the main file.

Figure 7: Simplified version of annotation tutorial.

[illegible]

A.3 Comparisons with Prior Egocentric VQA Benchmarks.

EgoMemoria [27], HourVideo [28], and EgoLifeQA [1], listing their lighting conditions (mainly daytime or nighttime), video duration length, the number of testing QA examples, number of QA type categories, if temporal-oriented tasks are included or emphasized, and the evaluation metric.

We highlight that EgoNight-VQA is the first to explore nighttime egocentric VQA with aligned day–night video pairs.

Dataset	Lighting	Video Length	# Test	# Categories	Temporal	Metric Type
EgoVQA	☀	(25s, 100s)	250	3	✗	OpenQA
EgoTaskQA	☀	25s	8k	4	✗	OpenQA
EgoSchema	☀	180s	500	-	✗	CloseQA
EgoThink	☀	-	750	12	✗	OpenQA
EgoTempo	☀	45s	500	10	✓	OpenQA
EgoCross	☀	23s	957	15	✓	CloseQA & OpenQA
EgoBlind	☀	(0s, 120s)	5311	6	✓	OpenQA
EgoMemoria	☀	(30s, 1h)	7026	-	✓	CloseQA
HourVideo	☀	(20min, 120min)	12976	-	✓	CloseQA
EgoLifeQA	mostly ☀	44.3 h	6000	5	✓	CloseQA
EgoNight-VQA	Aligned ☀ & ☾	(24s, 214s)	3658	12	✓	OpenQA

Table 5: Comparison between EgoNight-VQA and prior egocentric VQA benchmarks. ☀ means daytime, while ☾ indicates nighttime.

A.4 More EgoNight-VQA Explanations and Examples

A.4.1 QA Type Definition

We present the 12 QA types with their detailed definitions in Tab. 6.

QA Type	Attribute	Description
Object Recognition	Paired	Identify and recognize specific objects in the scene (e.g., “What is on the table?”).
Text Recognition	Paired	Read and interpret visible text or logos (e.g., “What does the sign say?”).
Spatial Reasoning	Paired	Understand spatial relations between objects (e.g., “What is left of the chair?”).
Scene Sequence	Paired	Recall the temporal order of visited scenes (e.g., “Which room did I enter after the kitchen?”).
Navigation	Paired	Working as an navigation assistant after watched the whole video (e.g., “How can I reach place B from place A?”).
Counting of Statics	Paired	Count static objects visible in the scene (e.g., “How many chairs are in the room?”).
Action Recognition	Paired	Identify human actions or interactions (e.g., “What action is being performed?”).
Non-Common-Sense Reasoning	Paired	Judge unusual or physically implausible cases, for synthetic videos. (e.g., “Is the door embedded inside the wall?”).
Lighting Recognition	Unpaired	Recognize the illumination source, also include counting. (e.g., “How many light sources are in the room?”).
Lighting Change	Unpaired	Detect changes in lighting conditions (e.g., “Did the light turn off during the clip?”).
Dynamic Detection	Unpaired	Detect dynamic moving objects (e.g., “Is a car/person moving across the scene?”).
Counting of Dynamics	Unpaired	Count the number of dynamic objects or events (e.g., “How many people walked by?”).

Table 6: **Detailed descriptions of the 12 QA types in EgoNight-VQA.** Paired QA types share the same QAs across day–night counterparts, while unpaired QA types are evaluated only at nighttime.

A.4.2 QA Examples

We show more QA examples of EgoNight-Synthetic, EgoNight-Sofia, EgoNight-Oxford in Fig. 9, Fig. 10, and Fig. 11, respectively. Note that for those paired QA types, we show both day and night frames, while for those unpaired QA types, we demonstrate nighttime frames only. Three frames are shown if the QA is spatial or static related, while more frames are given if the QA is temporal or more dynamic related.

A.5 More Experiments Setups

A.5.1 Setups for EgoNight-VQA Experiments.

In this section, we describe the model setup, how GPT is used as the judge, the prompting strategy, the GPU resources, and the approximate runtime for each dataset. For the closed-source model, we directly use the API call. For open source models, we use LLama-Factory [89] except VideoLLama3 [33] and EgoGPT [1]. We use NVIDIA A6000 GPUs for all the model, except for Qwen2.5-VL-72B, we use 2 NVIDIA H200 GPUs for larger GPU memory. The inference speed for each model is shown in Tab. 7.



Figure 9: More QA examples from EgoNight-Synthetic dataset.

Model	Inference Speed (min)
GPT-4.1	<5
Gemini 2.5 Pro	<5
InternVL3-8B	<5
Qwen2.5-VL-72B	25
Qwen2.5-VL-7B	<5
Qwen2.5-VL-3B	<5
GLM-4.1V-9B-Base	<5
VideoLLaMA3-7B	<5
LLaVA-NeXT-Video-7B	50
EgoGPT	<5

Table 7: Inference speed for different models (per Video).

Video frames are sampled at 2 fps for EgoNight-Synthetic, and 1 fps for EgoNight-Sofia and EgoNight-Oxford, without imposing a maximum frame limit. To further ensure fairness and consistency, the exact prompts used for each task are provided below.

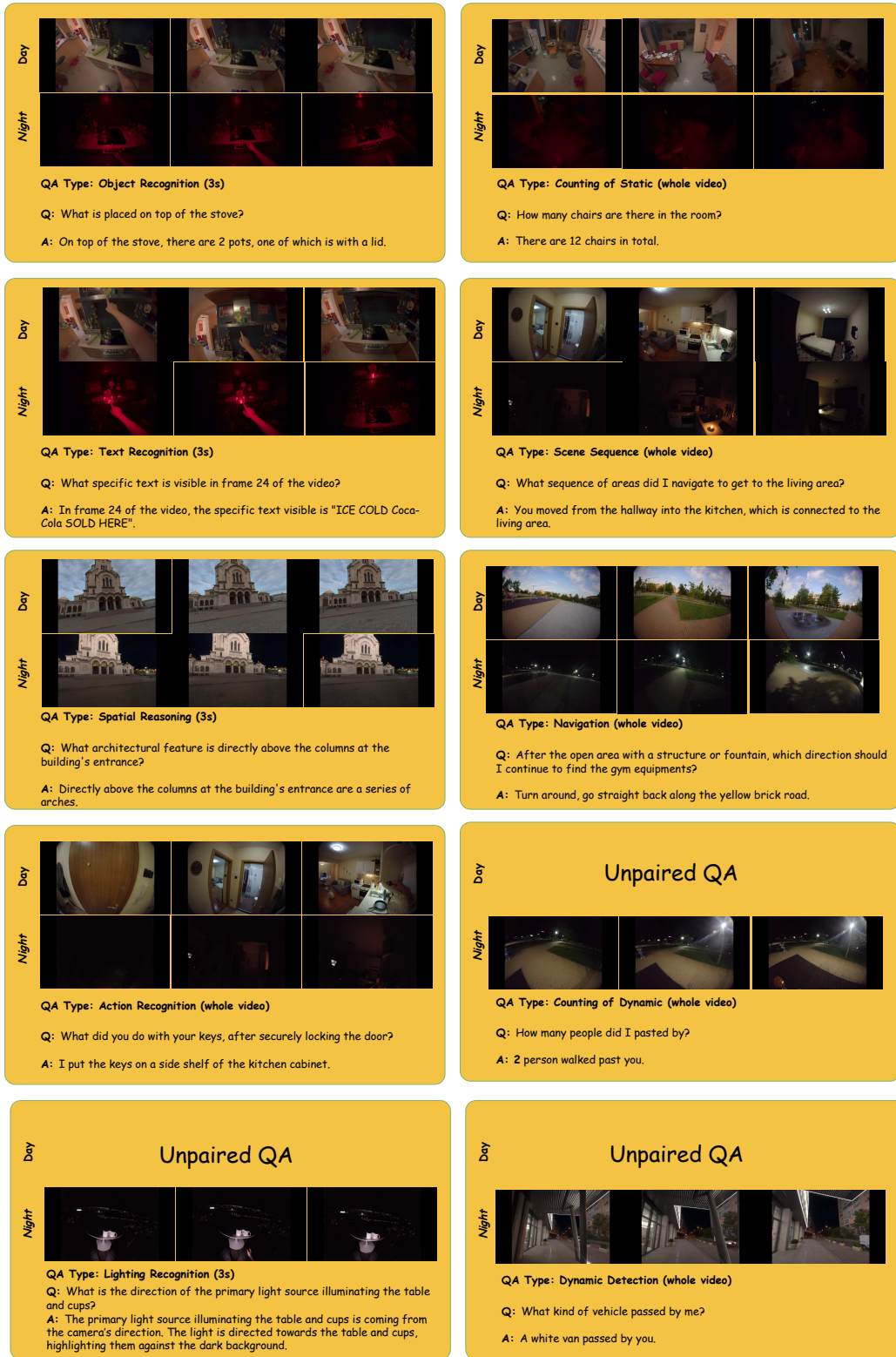


Figure 10: More QA examples from EgoNight-Sofia dataset.

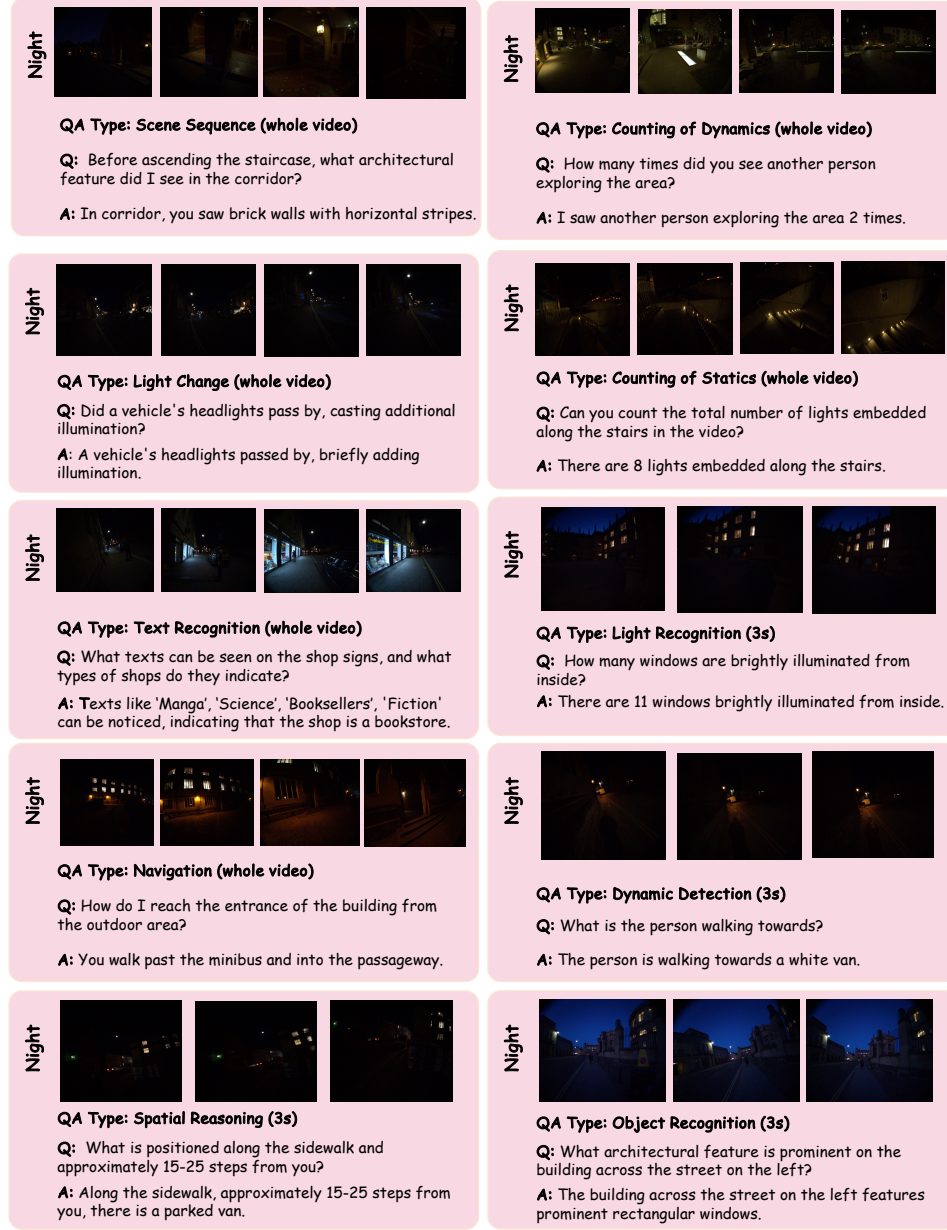


Figure 11: More QA examples from EgoNight-Oxford dataset.

Model Evaluation Prompt. For evaluating the language model, we use the following prompt:

Please carefully read the question, use the visual cues in the {video} to answer the question: {question}.

The original FPS of the video is {original_video_fps}. This image set is obtained by sampling at {sampling} fps.

Do not include any other content. You need to answer the question in any case and not demand additional context information.

Note: All the actions mentioned refer to the person who recorded the video.

Evaluation Protocol. Since the questions and answers are open-ended, we utilize GPT-4.1 [36] as a judge. Here is the prompt for evaluating the score given the model prediction, the ground truth answer, and the corresponding question:

role: system,
content: You are an intelligent chatbot designed for evaluating the correctness of AI assistant predictions for question-answer pairs.
Your task is to compare the predicted answer with the ground-truth answer and determine if the predicted answer is correct or not. Here's how you can accomplish the task:
INSTRUCTIONS:
1. Focus on the correctness and accuracy of the predicted answer with the ground-truth.
2. Consider uncertain predictions, such as 'it is impossible to answer the question from the video', as incorrect, unless the ground truth answer also says that.
role: user,
content: Please evaluate the following video-based question-answer pair:
Question: {question}
Ground truth correct Answer: {answer}
Predicted Answer: {predicted_answer}
Provide your evaluation as a correct/incorrect prediction along with the score where the score is an integer value between 0 (fully wrong) and 5 (fully correct). The middle score provides the percentage of correctness. For question that counting the number of objects, if the predicted answer falls in the range of the ground truth answer, it should be considered as correct.
Please generate the response in the form of a Python dictionary string with keys 'pred', 'score' and 'reason', where value of 'pred' is a string of 'correct' or 'incorrect', value of 'score' is in INTEGER, not STRING and value of 'reason' should provide the reason behind the decision."

A.5.2 Setups for Day-Night Correspondence Retrieval.

In this section, we describe the setup of the model, method, e.g. how feature-based retrieval for vision encoders, how prompt VLMs, metric, result, GPU, cost time, and other details.

i) Spatial Retrieval (Place Recognition). For feature-based methods [84, 85], we calculate the CLS tokens f^i of each frame within the video clip with the vision encoder, with i the frame index in the clip. Then, the "best matching" strategy is implemented to calculate the similarity between the query clip v_q and the database clip v_d . The best cosine similarity between the features of the query clip f_q^i and the database clip f_d^j ,

$$\sigma(v_q, v_d) = \max_{i \in [0, s-1], j \in [0, s-1]} \cos(f_q^i, f_d^j). \quad (1)$$

The database video clips are ordered based on the similarity σ and then the most similar clip is retrieved. Similarly, for MLLM-based methods [36, 32], we ask the MLLM to assess the "pairwise" similarity between each query-database pair and order the database clips by similarity. The prompt to the MLLM is as follows:

You are given two video clips from different scenes.
Your task is to evaluate how similar these two scenes are based on their spatial layout, furniture, objects, architectural features, and overall room structure.

CLIP STRUCTURE:
– Images 1–>(s–1) from Query Scene
– Images s–>(2s–1) from Database Scene

TASK:
Please carefully analyze and compare the spatial layout, furniture placement, objects, architectural features, and overall room structure between these two video clips.

IMPORTANT: Please respond with ONLY a single numerical similarity score between 0.0 and 1.0, where:
– 0.0 = Completely different scenes (different rooms/locations)
– 1.0 = Identical or nearly identical scenes (same room/location)
– Values in between represent varying degrees of similarity

Example responses: "0.85", "0.23", "0.67"
1.0 should be used when the two scenes are identical, so don't use 1.0 if the two scenes are not 100% identical.
Please provide only the numerical score without any additional text or explanation.

It is noticeable that existing MLLMs have difficulty in processing long-horizon and multi-scene videos. We also conduct the “all-in-one-prompt” experiments by inputting all the images of the query clip and the database clips in one prompt and asking the MLLM to output the ordered database clips. The “all-in-one-prompt” strategy leads to largely degraded performance, as shown in Tab. 8.

Prompt Strategy	Spatial Retrieval R@1 - Synthetic	
	Day $\rightarrow$ Day	Night $\rightarrow$ Day
Pairwise	75.6	54.1
All-in-one	10.5	28.5

Table 8: Ablation on prompting for night-to-day spatial retrieval task.

ii) Temporal Localization. The mIoU metric is defined as:

$$\text{mIoU} = \frac{1}{M} \sum_{m=1}^M \frac{|[t_i, t_j] \cap [t_i^*, t_j^*]|}{|[t_i, t_j] \cup [t_i^*, t_j^*]|}, \quad (2)$$

where M denotes the total number of meta-tasks (1000 in our setup). For feature-based temporal localization, we apply the “best-match” strategy similar as spatial localization, localizing the query clip to the frame stamp with the best clip-to-clip similarity:

$$i = \arg \max_i \sigma(v_q, v_d), v_d = v_D[i : i + s], \quad (3)$$

where v_D is the parent full video and the end frame will be $i + s - 1$. For MLLM-based method, we input the query clip and the parent full video in the prompt and ask the MLLM to output the start and end frame of the query within the full video. The prompt is as follows:

*You are given a query video clip and a complete video sequence from the same scene.
Your task is to find the exact temporal position where the query clip appears in the complete video sequence.*

IMPORTANT CONTEXT:

- The query clip shows s consecutive frames from a video sequence
- The complete video sequence shows ALL frames from the same scene in chronological order
- The query clip appears as a consecutive subsequence somewhere within the complete video sequence
- You need to find the exact start and end frame numbers where this subsequence appears

IMAGE STRUCTURE:

- Images $1 \rightarrow s$: Query video clip (consecutive frames to find)
- Images $\{s+1\} - \{s+1+\text{video_len}\}$: Complete video sequence (all frames in chronological order)

TOTAL IMAGES: $\{\text{query_count} + \text{database_count}\}$ images

TASK:

1. Look at the query clip to understand what sequence you’re looking for
2. Search through the complete video sequence to find where this exact sequence appears
3. The query sequence should appear as consecutive frames in the complete video sequence
4. Pay attention to camera movements, object positions, and scene changes to identify the matching sequence

FRAME NUMBERING:

- The complete video sequence frames are numbered from $\{\min(\text{database_frame_numbers})\}$ to $\{\max(\text{database_frame_numbers})\}$
- You need to return the actual frame numbers from this range

RESPONSE FORMAT:

Respond with **ONLY** two numbers separated by a comma: “start_frame,end_frame”

- start_frame: The frame number where the query clip begins in the complete video sequence
- end_frame: The frame number where the query clip ends in the complete video sequence

Example: If the query clip appears at frames 15–19 in the complete sequence, respond: “15,19”

Valid frame range: $\{\min(\text{database_frame_numbers})\}$ to $\{\max(\text{database_frame_numbers})\}$

A.5.3 Setups for Egocentric Depth Estimation at Night.

We evaluate four off-the-shelf monocular depth systems without night-specific fine-tuning. For each, we highlight features pertinent to our setting. (F) denotes support for fisheye egocentric images; (U) denotes undistorted/pinhole images.

1. **Depth Anything V2 (metric).** (U) Foundation MDE model (DPT head with DINOv2 backbone) trained on large-scale synthetic labels plus pseudo-labeled real images. We use the *official metric* checkpoints: *Indoor* (Hypersim-tuned) for indoor frames and *Outdoor* (VKITTI2-tuned) for outdoor frames. Outputs metric depth in meters and is known for strong zero-shot generalization.
2. **StreamVGGT.** (U) A causal/streaming transformer for video geometry that processes frames sequentially with state caching to improve temporal consistency and enable real-time inference. We run it in streaming mode to obtain per-frame depth on egocentric sequences.
3. **Depth Any Camera (DAC).** (F) Zero-shot *metric* depth across diverse camera models via a unified ERP (equirectangular) representation with pitch-aware image-to-ERP conversion and FoV alignment. We use the official release with default settings on our pinhole inputs.
4. **UniK3D.** (F) Universal-camera monocular 3D estimation with a spherical 3D formulation and a learned “pencil-of-rays” camera module, enabling accurate metric depth across pinhole, fisheye, and panoramic views. We run the official model in eval mode; when available, we provide intrinsics for pinhole frames.

A.6 More Experimental Results

In this section, we show models per QA accuracy on each dataset for EgoNight-Synthetic in Tab. 9, EgoNight-Sofia in Tab. 10, and EgoNight-Oxford in Tab. 11.

Model	Object Rec.	Text Rec.	Spatial	Scene Seq.	Nav.	Light Rec.	Cnt. Static	Non-Common	Avg.
<i>Closed-Source MLLMs</i>									
Gemini	25.94	39.39	32.43	35.47	30.77	31.97	21.88	15.15	28.34
GPT-4.1	25.19	54.55	35.42	28.44	27.09	35.25	20.83	16.67	27.75
<i>Open-Source MLLMs</i>									
InternVL3-8B	17.29	10.61	28.34	20.80	10.37	18.03	16.93	21.21	18.97
Qwen2.5-VL-72B	15.41	16.67	28.88	21.41	7.02	12.30	10.94	21.21	17.15
Qwen2.5-VL-7B	6.77	13.64	17.98	11.93	9.03	15.57	14.06	18.69	13.26
Qwen2.5-VL-3B	7.89	22.73	17.17	14.37	11.37	8.20	13.02	12.63	13.06
GLM-4.1V-9B-Base	13.16	36.36	23.71	21.71	7.02	15.57	19.27	14.14	17.69
LLaVA-NeXT-Video-7B	5.26	10.61	10.08	4.59	3.01	16.39	4.69	9.60	6.85
VideoLLaMA3-7B	10.90	21.21	19.07	23.55	7.02	7.38	18.49	16.67	15.97
<i>Egocentric MLLMs</i>									
EgoGPT	6.02	19.70	18.53	19.88	8.36	8.20	17.71	17.68	14.79
<i>Average across all models</i>									
Average	13.38	24.55	23.16	20.21	12.11	16.89	15.78	16.36	17.38

Table 9: Night-time VQA accuracy (%) per model across all QA categories for EgoNight-Synthetic.

Model	Object Rec.	Text Rec.	Spatial	Scene Seq.	Action	Nav.	Light Rec.	Cnt. Static	Light Dyn.	Dynamic	Cnt. Dynamic	Avg.
<i>Closed-Source MLLMs</i>												
GPT-4.1	24.44	33.78	41.32	30.09	38.10	27.27	38.33	24.62	30.00	13.33	25.00	31.06
Gemini	32.22	47.30	35.54	24.78	34.92	29.29	43.33	27.69	15.00	26.67	40.00	32.67
<i>Open-Source MLLMs</i>												
InternVL3-8B	16.67	17.57	33.88	24.78	22.22	21.21	20.00	22.31	25.00	6.67	15.00	22.61
Qwen2.5-VL-72B	14.44	21.62	36.36	18.58	19.05	25.25	18.33	13.85	20.00	6.67	20.00	20.99
Qwen2.5-VL-7B	7.78	10.81	21.49	18.58	11.11	18.18	1.52	15.38	9.52	6.25	15.00	14.02
Qwen2.5-VL-3B	11.11	8.11	23.97	13.27	11.11	16.16	10.00	15.38	5.00	13.33	10.00	14.16
GLM-4.1V-9B-Base	12.22	12.16	30.58	15.04	14.29	17.17	11.67	25.38	15.00	6.67	10.00	18.14
VideoLLaMA3-7B	3.33	5.41	13.22	11.50	6.35	9.09	8.33	18.46	10.00	20.00	15.00	10.68
LLaVA-NeXT-Video-7B	8.89	5.41	18.18	12.39	12.70	15.15	11.67	13.85	0.00	13.33	20.00	12.67
<i>Egocentric MLLMs</i>												
EgoGPT	9.18	3.41	19.26	16.54	6.67	7.96	10.00	14.97	0.00	0.00	10.00	11.47
<i>Average across all models</i>												
Average	13.99	16.31	27.29	18.53	17.45	18.53	17.16	19.13	12.94	11.26	18.00	18.76

Table 10: Night-time VQA accuracy (%) per model across all QA categories for EgoNight-Sofia.

Models	Object Rec.	Text Rec.	Spatial	Scene Seq.	Action	Nav.	Light Rec.	Cnt. Static	Light Dyn.	Dynamic	Cnt. Dynamic	Avg.
<i>Closed-Source MLLMs</i>												
GPT-4.1	64.52	34.88	41.35	34.13	65.59	43.75	30.43	18.49	18.52	37.84	18.52	38.95
Gemini 2.5 Pro	56.14	46.51	38.30	27.27	59.55	32.65	31.11	18.97	23.33	37.14	3.70	34.83
<i>Open-source MLLMs</i>												
InternVL3-8B	40.00	27.91	18.68	14.66	36.78	20.83	6.98	9.91	6.67	37.14	7.41	20.57
Qwen2.5-VL-72B	38.18	39.53	29.67	13.79	22.99	23.96	16.28	17.12	16.67	11.43	11.11	22.07
Qwen2.5-VL-7B	9.09	23.26	20.88	11.21	10.34	15.62	9.30	11.71	3.33	11.43	18.52	13.35
Qwen2.5-VL-3B	14.55	13.95	9.89	13.79	19.54	20.83	6.67	10.81	3.70	17.14	6.98	13.62
GLM-4V	22.81	30.23	21.43	15.52	28.74	9.38	19.57	17.12	6.67	37.14	14.81	19.57
VideoLLaMA3-7B	20.00	11.63	15.38	9.57	10.34	7.29	9.52	9.01	0.00	17.65	7.41	10.81
LLaVA-NeXT-Video-7B	9.09	4.65	10.99	0.00	0.00	0.00	6.98	0.00	0.00	2.86	0.00	2.86
<i>Egocentric MLLMs</i>												
EgoGPT	21.13	27.08	14.14	3.42	14.61	6.25	11.11	6.25	21.62	18.92	17.86	12.44
<i>Average across all models</i>												
Avg.	29.81	25.98	22.32	14.55	27.16	18.09	14.96	12.01	10.75	22.95	10.33	19.04

Table 11: Night-time VQA accuracy (%) per model across all QA categories for EgoNight-Oxford.

A.7 More Visualization Results

A.7.1 EgoNight-VQA

Here we provide more examples of comparison between day & night VQA output in Fig. 12, which clearly shows the gap between day and night video understanding.

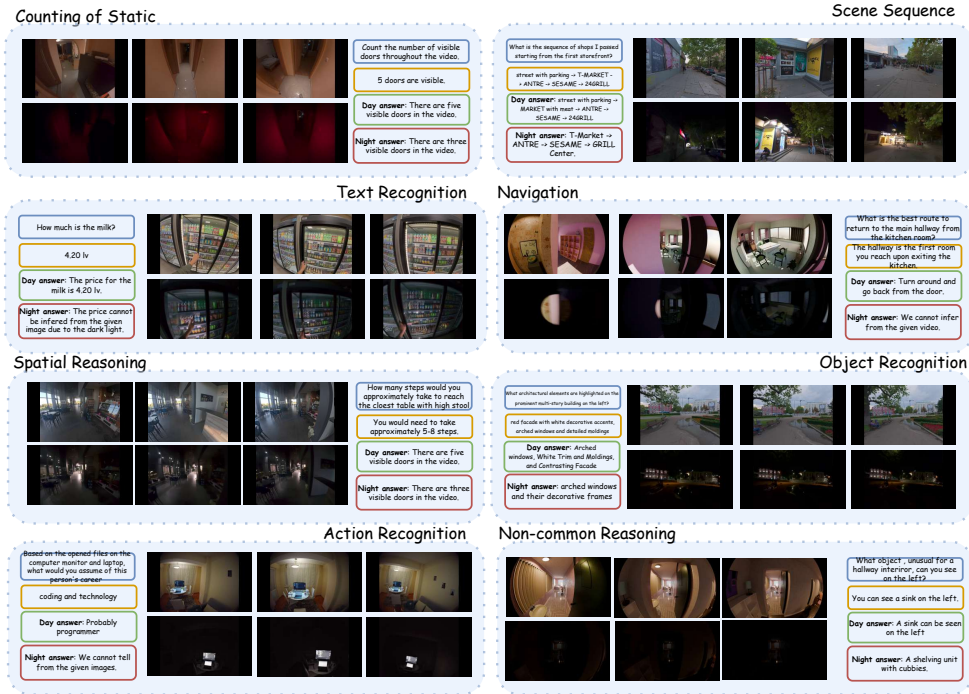


Figure 12: More QA examples with day and night answer produced by the same model.

A.7.2 Day-Night Correspondence Retrieval

We visualize the qualitative result on one meta sample of Night-to-Day spatial retrieval to better demonstrate the experiment setup and the performance of the benchmarked methods. As shown in Fig. 13, the light condition of the query video clip is drastically different from that of the database clips. Such a difference imposes a great challenge for existing methods in distinguishing the target scenes from the other candidate databases' clips, showing the value of the dataset in the place recognition task.

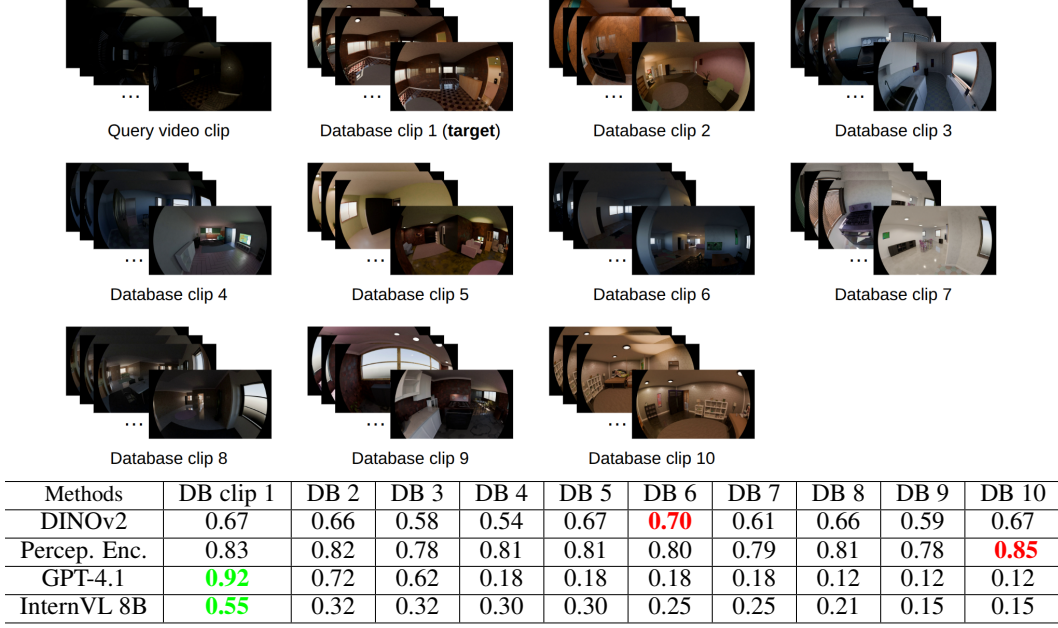


Figure 13: Qualitative Result on one meta sample of spatial retrieval. The query video clip and the database video clips are visualized in the image. The table below the figure shows the similarity score between the query and the database clips calculated with different methods. The most similar one is in **bold**, and correct retrieval is in **green**, and the incorrect one is in **red**.

A.7.3 Egocentric Depth Estimation at Night

We provide additional qualitative results across day–night conditions (Figs. 14, 15, 16, 17). Consistent with the main paper, nighttime is substantially more challenging: low SNR, head-motion blur, extreme dynamic range, color/white-balance shifts, and auto-exposure fluctuations amplify scale ambiguity and erode edge fidelity, leading to over-smoothed surfaces, depth collapse in dark regions, halos around bright point sources, and temporal instability. UniK3D remains the strongest overall in preserving scene structure under these conditions, though performance still degrades under extreme darkness and sparse texture. By contrast, *StreamVGGT* and *DAC* are notably brittle at night, frequently washing out structure, misinterpreting specular highlights, and producing flattened or unstable depth in large low-illumination areas. The effect is most pronounced outdoors in EgoNight-Sofia and EgoNight-Oxford, where wide dynamic range, sparse texture, and point-light saturation further depress accuracy across methods.

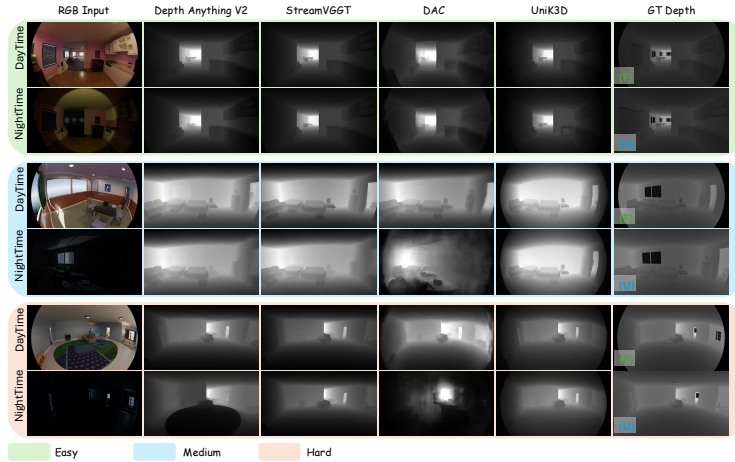


Figure 14: Qualitative results of monodepth estimation in day and night on EgoNight-Synthetic dataset according to different difficulty levels.

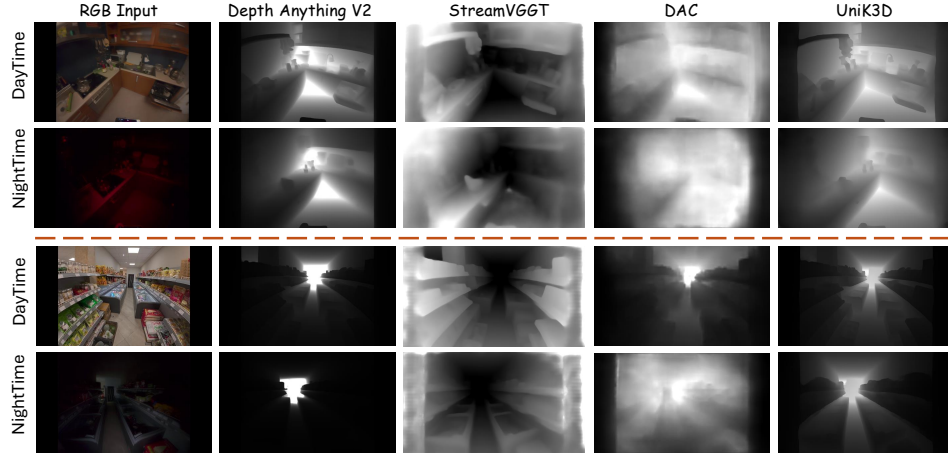


Figure 15: Qualitative results of monodepth estimation in day and night on EgoNight-Sofia, indoor.

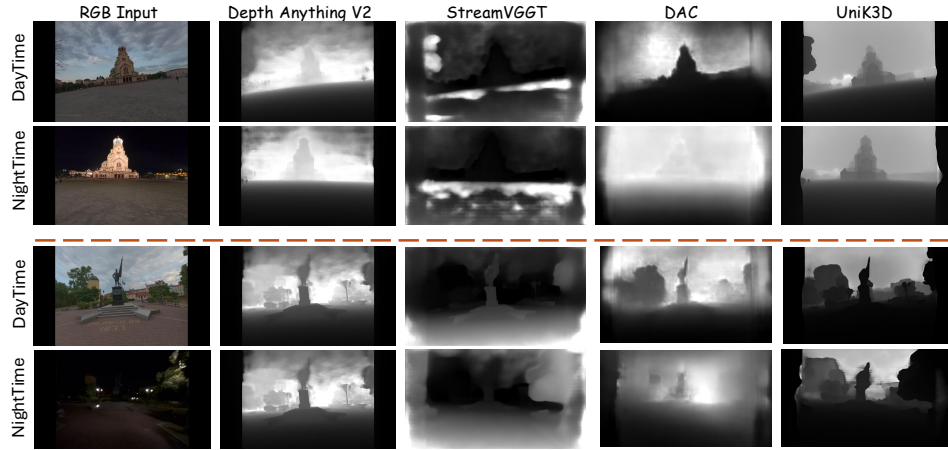


Figure 16: Qualitative results of monodepth estimation in day and night on EgoNight-Sofia, outdoor.

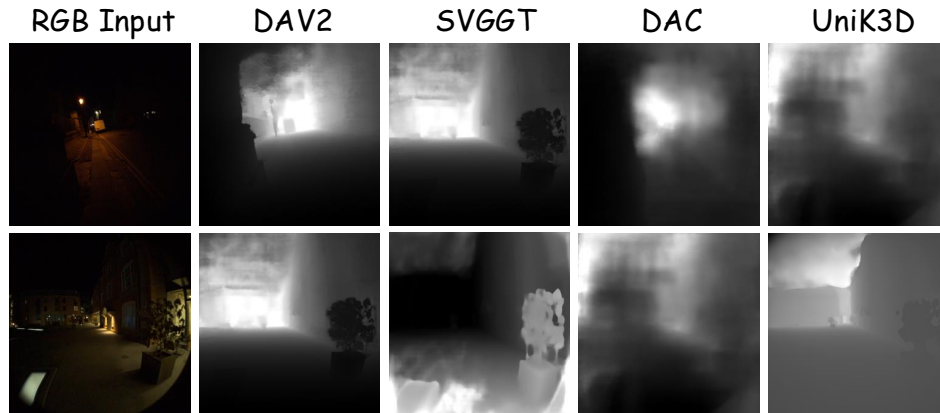


Figure 17: Qualitative results of monodepth estimation in day and night on EgoNight-Oxford dataset, note that DAV2 and SVGGT are shortened for Depth Anything V2 and StreamVGGT, respectively.

A.8 More Analysis

A.8.1 Limitations

We acknowledge two main limitations of EgoNight. (1) The dataset scale remains modest compared to large-scale vision–language corpora. However, as a testbed, we argue that the current scale of 3,600+ human-verified QA pairs is already sufficient for benchmarking. In future work, we plan to further scale up nighttime videos by synthesizing more data and recording additional real-world footage, which will enable not only benchmarking but also pretraining and fine-tuning to improve MLLM performance. (2) EgoNight primarily focuses on day–night illumination shifts, while other real-world challenges such as weather variations (rain, fog) and extreme camera motion are not covered. We view these as promising directions for future extensions of EgoNight.

A.8.2 Contribution to the Community

We believe EgoNight will serve as a valuable resource for the research community in several ways. First, it provides the *first benchmark suite* dedicated to egocentric nighttime vision, a long-overlooked but practically critical setting for robust AI assistants. Second, the dataset’s unique day–night alignment enables rigorous analysis of illumination effects, offering insights that cannot be obtained from prior egocentric benchmarks. Third, by covering multiple tasks, VQA, day-night correspondence retrieval, and depth estimation, EgoNight provides a comprehensive testbed that can catalyze progress across both perception and reasoning. Finally, with all data, annotations, and evaluation code to be released publicly, EgoNight is designed to be easily accessible, extensible, and reproducible, supporting future research on egocentric vision understanding learning.

A.8.3 Usage of Large Language Models (LLMs)

Our annotation pipeline and benchmark evaluation both leverage large language models (LLMs). For data construction, advanced multimodal LLMs are used to generate initial captions, questions, and pseudo answers, which are then refined by human annotators. This hybrid model–human approach substantially reduces annotation cost while ensuring quality. For evaluation, we adopt the *LLM-as-a-Judge* paradigm to assess the semantic correctness of model outputs against ground-truth answers, following recent practice in egocentric VQA. Beyond annotation and evaluation, we also used LLMs to support paper preparation, such as generating icons for illustration figures and assisting with proof-reading. Importantly, while LLMs serve as practical tools throughout our workflow, all core ideas, dataset design, experiments, and analyses are conceived and conducted independently by the authors.