

Reinforcement Learning Tuning for VideoLLMs: Reward Design and Data Efficiency

Hongyu Li^{1*}, Songhao Han^{1*}, Yue Liao^{2*}, Junfeng Luo³, Jialin Gao³,
Shuicheng Yan², Si Liu^{1‡}

¹ Beihang University ² National University of Singapore ³ Meituan

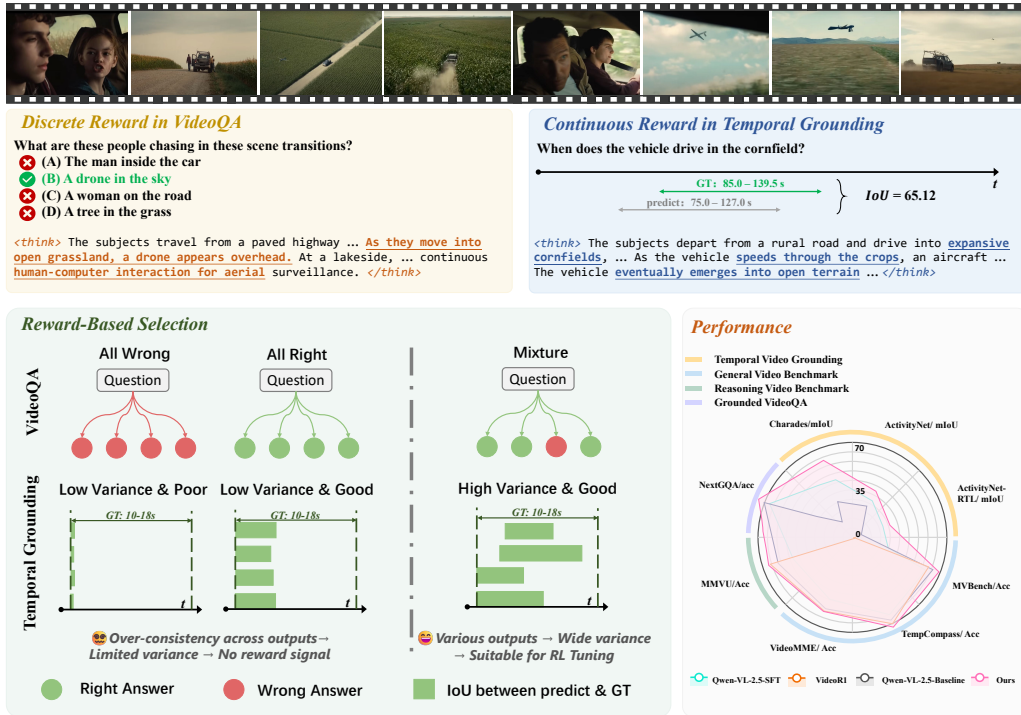


Figure 1: Overview of our reward-based framework for post-training on video understanding tasks. *Top left:* Semantic reasoning is supervised using a discrete reward from multi-choice VideoQA, where the model selects the correct answer. *Top right:* Temporal reasoning is guided by a continuous reward based on temporal IoU in grounding tasks. *Bottom left:* A reward-based data selection strategy filters training samples based on behavioral variance. High-variance samples induce stronger preference gradients and are prioritized for GRPO optimization. *Bottom right:* Performance comparison across baselines and our Temporal-RLT, showing consistent improvements on QA, grounding, and reasoning tasks.

Abstract

Understanding real-world videos with complex semantics and long temporal dependencies remains a fundamental challenge in computer vision. Recent progress in multimodal large language models (MLLMs) has demonstrated strong capabilities in vision-language tasks, while reinforcement learning tuning (RLT) has further improved their reasoning abilities. In this work, we explore RLT as a post-training

*Equal contribution.

†Project leader.

‡Corresponding author.

strategy to enhance the video-specific reasoning capabilities of MLLMs. Built upon the Group Relative Policy Optimization (GRPO) framework, we propose a dual-reward formulation that supervises both semantic and temporal reasoning through discrete and continuous reward signals. To facilitate effective preference-based optimization, we introduce a variance-aware data selection strategy based on repeated inference to identify samples that provide informative learning signals. We evaluate our approach across eight representative video understanding tasks, including VideoQA, Temporal Video Grounding, and Grounded VideoQA. Our method consistently outperforms supervised fine-tuning and existing RLT baselines, achieving superior performance with significantly less training data. These results underscore the importance of reward design and data selection in advancing reasoning-centric video understanding with MLLMs. Our code is available at [Temporal-RLT](#).

1 Introduction

Video is a fundamental medium for capturing the dynamics of the visual world. Understanding real-world videos—characterized by complex semantics, diverse visual content, and long-range temporal dependencies—remains a core challenge in computer vision. While large multimodal models (MLLMs) have shown promising capabilities in vision-language tasks, recent efforts have begun extending them to video understanding [3, 30, 2, 14, 39, 15], where temporal alignment and structured reasoning are critical. Inspired by human cognitive processes, incorporating explicit reasoning into large language and multimodal models has been shown to significantly enhance their comprehension, particularly in tasks requiring multi-step inference and temporal understanding. Reinforcement learning tuning (RLT) [7, 26, 22, 4, 27] has emerged as an effective post-training paradigm for stimulating such reasoning abilities, with notable success in both language and vision-language domains through preference-based optimization.

In this work, we propose **Temporal-RLT**, a RLT framework designed to improve the video-specific reasoning capabilities of MLLMs. Our method builds upon the Group Relative Policy Optimization (GRPO) algorithm [26], which optimizes model behavior by comparing multiple sampled outputs and applying fine-grained, verifiable reward signals. While GRPO has been successfully applied in textual reasoning tasks, its extension to video domains remains largely underexplored and introduces new challenges in modeling structured temporal reasoning under multimodal supervision.

The first challenge is *"how to design effective reward functions that incentivize video-specific reasoning."* To address this, as shown in Fig. 1, we propose a dual-reward formulation that combines discrete rewards from multi-choice VideoQA with continuous rewards derived from temporal Intersection over Union (tIoU) in temporal grounding tasks. The discrete reward promotes semantic alignment by supervising answer correctness, while the continuous reward guides temporal localization accuracy. We further extend this formulation to Grounded VideoQA, where the model must both answer questions and localize supporting video segments, thereby enhancing interpretability and temporal grounding.

The second challenge involves identifying *"effective training samples for GRPO-based RLT"*. Since GRPO relies on behavioral variance among sampled outputs to generate preference gradients, we propose a variance-aware data selection strategy based on repeated inference. For each sample, we perform multiple forward passes using a base VideoLLM and measure prediction variability. In multi-choice QA, we retain samples with moderate answer consistency to ensure a mix of correct and incorrect responses. In temporal grounding, we compute the gap between maximum and average IoU to identify samples with high intra-group variance. This approach filters out uninformative examples and focuses training on samples most likely to produce meaningful learning signals under GRPO.

Based on this strategy, we construct two datasets to support Temporal-RLT. **Temporal-RLT-Full-490k** serves as a comprehensive dataset covering VideoQA, Temporal Grounding, and Grounded VideoQA. From this data pool, we derive **Temporal-RLT-32k**, a high-quality subset selected for its diversity and optimization potential, enabling efficient training with significantly less data.

We validate Temporal-RLT across 8 benchmarks spanning General VideoQA, Reasoning VideoQA, Temporal Grounding, and Grounded VideoQA. Our method consistently outperforms strong supervised fine-tuning (SFT) baselines, and achieves superior performance compared to Video-R1 [4].

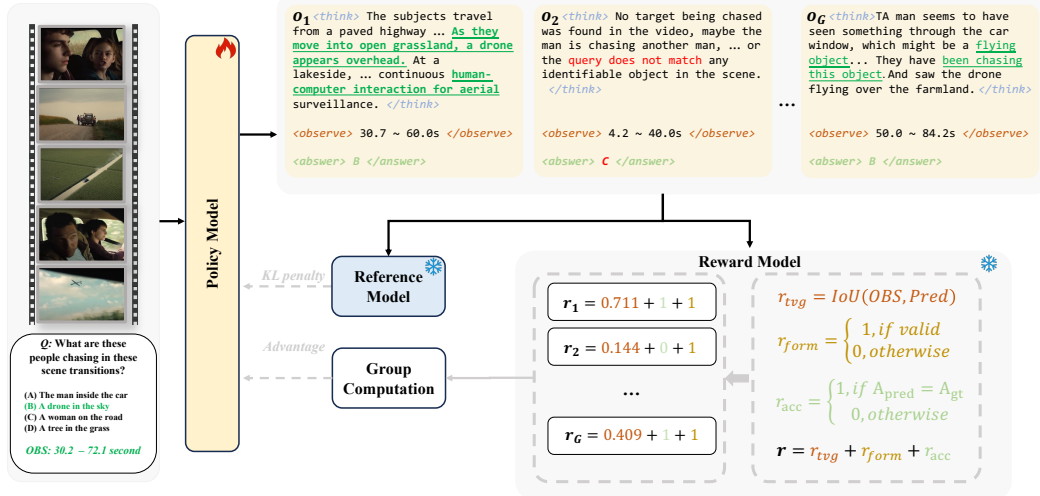


Figure 2: Overview of our training framework. Given a video and its corresponding instruction, the VideoLLM generates either an answer choice for multi-choice QA or a temporal segment for grounding. The model is optimized using a GRPO-based RLT pipeline with task-specific rewards, including a discrete reward for QA and a continuous temporal IoU reward for grounding.

Additionally, we conducted extensive ablation studies to verify the effectiveness of our data selection strategy and the performance gains achieved by various rewards across different tasks. These results highlight the effectiveness of our reward design and data selection strategy for enhancing video-specific reasoning, while demonstrating strong performance with data efficiency in RLT-based video understanding.

2 Temporal-RLT

In this section, we present *Temporal-RLT*, our reinforcement learning tuning pipeline tailored for video large language models (VideoLLMs). As illustrated in Fig. 2, Temporal-RLT extends a standard VideoLLM architecture by integrating a GRPO-based RLT module to further enhance the model’s video-specific reasoning and comprehension capabilities. We begin by outlining the overall pipeline of the base VideoLLM (Sec. 2.1). We then introduce our dual reward formulation (Sec. 2.2), consisting of a discrete semantic reward from multi-choice QA (MC-QA) and a continuous temporal reward based on temporal IoU (tIoU). Finally, we describe how GRPO is incorporated into the Temporal-RLT framework under this dual-reward setting, forming a unified post-training strategy for video understanding (Sec. 2.3).

2.1 Base Model: QwenVL-2.5 with Structured Reasoning

We adopt QwenVL-2.5 [2] as the base VideoLLM in our framework. QwenVL-2.5 is a powerful open-source multimodal language model pre-trained on large-scale video-language data, demonstrating strong performance across various video understanding tasks. Given multimodal input—comprising video frames and a natural language instruction—the model autoregressively generates outputs such as answer choices or grounding spans in a textual format. To move beyond direct answer prediction and promote explicit reasoning, we introduce a structured output format that encourages the model to *think before answering*. Specifically, we define a standardized response template using `<think>` and `<answer>` tags:

`<think> reasoning trace </think> <answer> final response </answer>.`

This format enforces the model to articulate its reasoning process prior to providing the final answer, enabling more interpretable and consistent behavior during inference. To further incentivize structural adherence and output consistency, we introduce a format alignment reward R_{format} , which encourages the model to generate outputs that conform to the predefined structure.

2.2 Dual Verifiable Rewards in Video Tasks

Reward design plays a critical role in guiding optimization within the GRPO-based RLT framework. To enhance both global semantic and fine-grain temporal reasoning in VideoLLMs, we introduce two complementary reward signals. A discrete semantic reward, derived from multi-choice QA tasks, encourages global semantic alignment. A continuous temporal reward, based on temporal IoU from grounding tasks, promotes fine-grained temporal comprehension. Together, these rewards provide structured supervision across multiple reasoning dimensions.

Global Discrete Semantic Reward. To incentivize high-level semantic reasoning, we leverage the Multi-Choice Video Question Answering (MC-VQA) task. In this setting, the model is given a video V , a question Q , and a set of candidate answers $\{A_1, A_2, \dots, A_n\}$. The objective is to select the correct answer A_c , where c denotes the index of the ground truth option. This transforms open-ended reasoning into a classification problem over a constrained semantic space.

To provide a verifiable reward signal, we adopt a binary accuracy-based reward function:

$$R_{\text{acc}} = \begin{cases} 1, & \text{if } A_{\text{pred}} = A_{\text{gt}} \\ 0, & \text{otherwise} \end{cases} \quad (1)$$

where A_{pred} is the model’s selected answer and A_{gt} is the ground truth. This discrete reward guides the model to align its global semantic understanding with task-specific supervision.

Fine-Grained Temporal IoU Reward. To capture fine-grained temporal reasoning, we incorporate the Temporal Video Grounding (TVG) task, where the model is required to localize an event within the video timeline. Given a video V and a language query, the model predicts a temporal segment defined by start and end timestamps (S_p, E_p) , which is then compared against the ground truth interval (S_g, E_g) .

We define the reward as the temporal IoU between the predicted and ground truth segments:

$$R_{\text{IoU}} = \frac{\max(0, \min(E_p, E_g) - \max(S_p, S_g))}{\max(E_p, E_g) - \min(S_p, S_g)} \quad (2)$$

This continuous reward quantifies the alignment between predicted and reference segments, encouraging precise temporal localization. It provides a smooth gradient signal that complements the binary supervision from semantic QA, enabling the model to reason across multiple granularity levels.

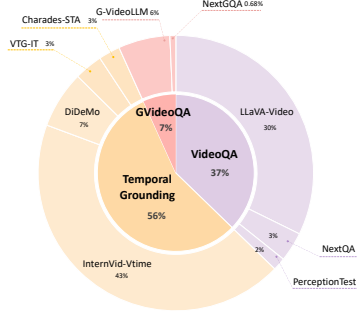
2.3 Temporal-RLT Optimization with GRPO

To optimize VideoLLMs under our task-specific reward formulations, we adopt GRPO, a preference-based reinforcement learning algorithm that fine-tunes models without requiring an explicit value function or learned critic. Instead of estimating absolute returns, GRPO compares the quality of multiple responses sampled for the same input and uses their relative rankings to guide optimization. This makes GRPO naturally compatible with our verifiable and multi-dimensional reward design.

We define three reward functions tailored to distinct video reasoning tasks. For Multi-Choice VideoQA (MC-QA), the reward is defined as $R_{\text{mc}} = R_{\text{format}} + R_{\text{acc}}$, where R_{acc} indicates answer correctness and R_{format} enforces output consistency. For Temporal Video Grounding (TVG), we use $R_{\text{tvG}} = R_{\text{format}} + R_{\text{IoU}}$, where R_{IoU} measures the alignment between predicted and ground-truth temporal segments. To bridge semantic and temporal reasoning, we further introduce Grounded VideoQA (GQA), where models are required to both answer the question and highlight relevant temporal spans using <observe> tags. The reward is defined as $R_{\text{gqa}} = R_{\text{format}} + \frac{1}{2}(R_{\text{acc}} + R_{\text{IoU}})$, encouraging joint semantic accuracy and temporal grounding.

For each input instance, a video-language instruction pair, the model samples a group of G candidate responses $\{o_1, o_2, \dots, o_G\}$. Each response is evaluated using the appropriate task-specific reward function, yielding a set of scalar rewards $\{R_1, R_2, \dots, R_G\}$. These rewards are then normalized within the group to compute the relative advantage for each sample:

$$A_i = \frac{R_i - \text{mean}(\{R_1, \dots, R_G\})}{\text{std}(\{R_1, \dots, R_G\})}, \quad (3)$$



Task	Data Source	Domain	Data Size
VideoQA	LLaVA-Video [14]	Diverse	158439
	NextQA [32]	Reasoning	17024
	PerceptionTest [23]	Perception, Reasoning	7392
Temporal Grounding	InternVid-Vtime [10]	Diverse	213158
	DiDeMo [1]	Diverse	33002
	VTG-IT [8]	Diverse	16318
	Charades-STA [6]	Indoor	12408
Grounded VideoQA	G-VideoLLM [29]	Vlog, News and Activity	29797
	NextGQA [33]	Reasoning	3358

Figure 3: Data statistic of Temporal-RLT-Full-490k training dataset. We find that converting all answers to an open-ended format is critical in reliably assessing question difficulty and effective model training.

where A_i reflects the relative quality of the i -th response. The model is updated to increase the likelihood of higher-scoring responses and reduce that of lower-scoring ones, enabling preference-driven fine-tuning guided by human-verifiable objectives.

By integrating our task-specific temporal rewards with GRPO’s relative optimization framework, *Temporal-RLT* provides a unified and scalable post-training strategy. It supports a wide range of video-language tasks, promotes structured and interpretable outputs, and eliminates the need for manually engineered reward critics.

3 Rethinking Data Efficiency for Video RLT

Data efficiency is critical for the post-training and fine-tuning of LLMs and MLLMs. It is equally essential in GRPO-style RLT for video understanding, where effective utilization of supervision is key to further improving model capability. In this section, we focus on the distinct characteristics of GRPO optimization and analyze data efficiency beyond conventional assumptions of data diversity. Building on our dual-reward formulation, we separately analyze the discrete and continuous reward paradigms, and propose a principled strategy to enhance reward-data utilization in the video domain.

3.1 Difficulty Estimation via Multi-Sampling

GRPO relies on relative preference signals within sampled groups to guide learning. Therefore, the informativeness of a training sample largely depends on whether its sampled outputs induce sufficient diversity in reward values. If all outputs are either highly correct or uniformly poor, the absence of contrast limits optimization signal. To quantify this, we adopt a unified multi-sampling-based difficulty estimation strategy. For each training instance, we perform $N \gg K$ inference runs using a base VideoLLM (pre-RLT) and compute reward scores across sampled outputs. For discrete tasks (e.g., multi-choice QA), we count the number of correct responses. For continuous tasks (e.g., temporal IoU), we evaluate the internal reward spread. Samples are then categorized based on their observed output variance, forming the basis for subsequent selection and prioritization.

3.2 Data Efficiency under Discrete Rewards

In multi-choice QA tasks, each sampled response receives a binary reward $r_i \in \{0, 1\}$ based on whether the predicted answer matches the ground truth. Effective GRPO optimization requires a *mixture* of correct and incorrect answers within the sampled group, so that lower-quality responses can be aligned toward higher-quality ones.

To approximate this mixture potential, we define difficulty based on the frequency of correct predictions among N base model inferences. We label a sample as:

$$\text{Easy if } c \geq \tau_{\text{easy}}, \quad \text{Hard if } c \leq \tau_{\text{hard}}, \quad \text{otherwise Medium,}$$

where c is the number of correct predictions out of N . We retain medium-difficulty samples to ensure sufficient behavioral variance, which supports more effective preference-based learning.

3.3 Data Efficiency under Continuous Rewards

In grounding tasks with continuous temporal IoU supervision, binary correctness no longer applies. Instead, we assess the spread of IoU scores within each sampled group to estimate difficulty. Intuitively, samples whose predictions yield a wide range of IoUs provide stronger relative learning signals. To avoid instability in direct variance computation under limited sampling (K), we define an approximate metric based on reward discrepancy:

$$\Delta_{\text{IoU}} = \max_i \text{IoU}_i - \text{mean}_i(\text{IoU}_i),$$

where a higher Δ_{IoU} indicates greater internal diversity. Samples with negligible spread are excluded from training, as they fail to generate preference gradients within the group.

3.4 Difficulty-Aware Data Construction

To support effective training under GRPO, we construct a difficulty-aware dataset for video RLT, comprising two specialized subsets: (1) Temporal-RLT-Full-490k as our foundational dataset containing diverse domain data for VideoQA, Temporal Video Grounding, and Grounded VideoQA tasks, and (2) Temporal-RLT-32k, a refined subset optimized for sample diversity and training efficacy. Our raw corpus includes a large set of video-question pairs and temporal grounding data with diverse semantics and temporal scopes, covering multiple benchmark datasets as visualized in Fig. 3.

We first conduct offline difficulty estimation on Temporal-RLT-Full-490k using the base VideoLLM through large-scale scaling inference. Each training instance is evaluated using our unified multi-sampling pipeline and assigned a task-specific difficulty score (based on accuracy frequency for discrete tasks and intra-group reward variation for continuous tasks). Based on this analysis, we derive Temporal-RLT-32k, a carefully curated training subset composed primarily of medium-difficulty samples that prove empirically more effective in generating stable and informative GRPO updates.

This difficulty-aware data construction improves sample efficiency by reducing redundant updates on uninformative (easy) samples and mitigating instability from noisy (hard) samples. It also ensures consistency across heterogeneous reward types, enabling a unified optimization framework for semantic and temporal reasoning tasks.

4 Experiment

In this section, we conduct comprehensive experiments across a range of General VideoQA, Temporal Grounding, and Video Reasoning benchmarks to evaluate the effectiveness of our Temporal-RLT framework. We further present an in-depth empirical analysis of data efficiency in GRPO-based RLT, highlighting the impact of our sample selection and dynamic data strategies.

4.1 Implementation Details

Training and Inference Details. All experiments are conducted using the Qwen-VL-2.5-7B [2] model, fine-tuned with GRPO-based reinforcement learning. We sample eight candidate responses per input during training, with a temperature of 1.0 and top- p of 0.99. Training is performed with a batch size of 48 on 8 NVIDIA A100 80GB GPUs. For inference, we adopt conservative decoding settings (temperature = 0.01, top- p = 0.001) to ensure output stability, following the official Qwen-VL-2.5 deployment configuration. Video inputs are processed at 2 FPS using Qwen-VL-2.5’s dynamic resolution module, with at least 4 tokens per frame and a maximum of 2048 tokens per video. Spatial resolution is automatically adjusted while preserving aspect ratio to ensure efficient computation. All training and inference procedures are conducted in BF16 precision for numerical stability.

Benchmarks. To evaluate our method across different aspects of video understanding, we conduct experiments on a diverse set of benchmarks. For general video-language understanding, we use MVBench [17], TempCompass [21], and VideoMME [5] (excluding subtitles). For video reasoning, we evaluate on MMVU [40], which focuses on multi-choice reasoning tasks. For temporal video grounding, we test on Charades-STA [6] and ActivityNet [13], while reasoning-driven temporal grounding is assessed on ActivityNet-RTL [11]. Finally, for grounded video question answering, we report results on the NextGQA test set [33]. These benchmarks collectively assess the model’s capabilities in semantic comprehension, temporal localization, and reasoning across diverse tasks.

Table 1: Main Experiment Results. Comparison of our Temporal-RLT framework against baseline VideoLLMs and post-training baselines across multiple video understanding and grounding benchmarks.

Method	Temporal Video Grounding			General VideoQA			Reasoning QA	Grounded QA	
	Charades	ANet	ANet-RTL	MVBench	TempCompass	VideoMME	MMVU	NextGQA	
	mIoU	mIoU	mIoU	Avg	Avg	Avg (wo sub)	Avg	mIoU	acc
General VideoLLM									
LLaMA-VID[18]	-	-	-	41.9	45.6	-	-	-	-
VideoLLaMA2[3]	-	-	-	54.6	-	47.9	44.8	-	-
LongVA-7B[39]	-	-	-	-	56.9	52.6	-	-	-
Video-UTR-7B[35]	-	-	-	58.8	59.7	52.6	-	-	-
LLaVA-OV-7B[14]	-	-	-	56.7	-	58.2	49.2	-	-
Kangaroo-7B[19]	-	-	-	61.1	62.5	56.0	-	-	-
GRPO-based Method and Baseline									
Qwen-VL-2.5[2]	28.0	24.0	6.0	65.3	70.9	56.1	61.3	20.2	77.2
Qwen-VL-2.5-SFT	43.0	24.3	18.1	62.0	68.7	49.6	52.5	28.3	70.6
Video-R1[4]	-	-	-	62.7	72.6	57.4	64.2	-	-
Temporal-RLT (ours)	57.0	39.0	27.6	68.1	73.3	57.6	65.0	37.3	78.7

4.2 Main Results

As shown in Tab. 1, our Temporal-RLT approach achieves significant improvements in multiple video understanding tasks compared to standard SFT. In temporal video grounding, our method yields substantial absolute gains: +14.0 mIoU on Charades, +14.7 on ActivityNet, and +9.5 on the reasoning-intensive ActivityNet-RTL benchmark. These results confirm that our RL-based strategy is highly effective at improving temporal localization. Importantly, while SFT often leads to performance degradation on general VideoQA benchmarks due to overfitting or misalignment between pretraining and downstream objectives, our method maintains strong generalization and consistently outperforms the SFT baseline across both general and reasoning-focused tasks. Specifically, we observe a +3.7 improvement on the MMVU benchmark, highlighting the model’s enhanced reasoning capability under our training paradigm. Moreover, when compared to Video-R1 [4], our method achieves superior performance across all evaluated tasks while requiring less data for training. This demonstrates the effectiveness of our data selection strategy and highlights the data efficiency of our RLT pipeline. Together, these results showcase the superiority of our Temporal-RLT framework in achieving strong performance across grounding, general QA, and reasoning tasks, while maintaining high data efficiency and strong generalization across diverse video benchmarks.

4.3 Ablation Study

We present a series of ablation studies to understand how task composition, reward design, and data selection affect the performance of our Temporal-RLT framework. The results reveal critical insights into optimization behavior and generalization across video-language tasks under GRPO-based tuning.

Multi-Task Ablation. Our experiments systematically evaluate the impact of different training data configurations on model performance. The TVG-only setup yields substantial improvements on temporal grounding tasks, with gains of +27.6 mIoU on Charades and +11.6 on ActivityNet, while maintaining reasoning performance (no degradation on MMVU) and improving grounded QA results (+5.8 mIoU). These findings suggest that enhanced temporal understanding can transfer effectively to related tasks. Adding VideoQA data introduces expected trade-offs: while it significantly improves question-answering accuracy on both general and reasoning-oriented VideoQA benchmarks, it leads to a moderate decline in grounding precision. Further augmenting with grounded VQA (G-VQA) produces more nuanced effects, including domain-specific gains on ActivityNet (+7.6) and smaller but meaningful improvements on Charades-STA (+1.1), along with a reasoning enhancement of +0.8 on MMVU. Together, these results highlight the partially decoupled development of temporal understanding and reasoning abilities. They also reveal critical trade-offs in data composition during RLT, illustrating the optimization boundaries across different video understanding tasks.

Performance with More Video Tokens. As shown in 3, by increasing the number of input video tokens, we achieve consistent performance improvements across four types of tasks. Moreover, compared to the SFT model, our method benefits more significantly from the increase in video token numbers. In contrast to the previous Video-R1, our approach demonstrates stable performance improvements across all tasks. Notably, when increasing the number of input video tokens, Video-R1

Table 2: Ablation Study on Multi-Task RL Tuning Components. We analyze and compare the effects of using different training data and reward types for Temporal-RLT.

Method	Temporal Video Grounding			General VideoQA			Reasoning QA	Grounded QA	
	Charades mIoU	ANet mIoU	ANet-RTL mIoU	MVBench Avg	TempCompass Avg	VideoMME Avg (wo sub)	MMVU Avg	NextGQA mIoU	acc
Qwen-VL-2.5 (baseline)	28.0	24.0	6.0	65.3	70.9	56.1	61.3	20.2	77.2
+w/ TVG	55.6	35.6	23.6	61.8	71.2	52.4	61.3	29.7	75.3
+w/ VQA	-	-	-	68.1	72.5	58.6	63.1	25.0	78.4
+w/ VQA + TVG	55.9	31.3	23.2	68.0	72.9	58.4	64.2	27.0	78.2
+w/ VQA + TVG + G-VQA	57.0	39.0	27.6	68.1	73.3	57.6	65.0	37.3	78.7

Table 3: Ablation Study on More Video Tokens. We analyze and compare the effects of using more Video Tokens for Inference Result .

Method	Video Tokens	Temporal Video Grounding			General VideoQA			Reasoning QA	Grounded QA	
		Charades mIoU	ANet mIoU	ANet-RTL mIoU	MVBench Avg	TempCompass Avg	VideoMME Avg (wo sub)	MMVU Avg	NextGQA mIoU	acc
Qwen-VL-2.5	2048	28.0	24.0	6.0	65.3	70.9	56.1	61.3	20.2	77.2
Qwen-VL-2.5	4096	35.7	33.4	7.1	66.5	71.0	59.4	61.1	26.3	76.5
Qwen-VL-2.5-SFT	2048	43.0	24.3	18.1	62.0	68.7	49.6	52.5	28.3	70.6
Qwen-VL-2.5-SFT	4096	43.8	23.9	18.7	63.0	68.9	52.1	55.2	27.4	71.3
Video-R1	2048	-	-	-	62.7	72.6	57.4	64.2	-	-
Video-R1	4096	-	-	-	63.9	73.2	59.3	63.8	-	-
Temporal-RLT (ours)	2048	57.0	39.0	27.6	68.1	73.3	57.6	65.0	37.3	78.7
Temporal-RLT (ours)	4096	58.0	41.3	33.0	69.0	73.6	59.6	66.1	38.2	79.7

experiences a performance drop on the reasoning QA task in the MMVU benchmark, whereas our method maintains robust gains.

Data Selection for Multi-choice VideoQA. The ablation study investigates the impact of difficulty distribution in multi-choice training data. As shown in Tab. 4a, reducing the proportion of easy samples while increasing medium-difficulty data improves both general and reasoning performance (evident from the comparison between the first and third rows). Similarly, increasing hard samples while reducing easy ones, while keeping medium data fixed, also enhances performance (first vs. second row), suggesting that easy data contribute little to effective training. Rows two through five further show that gradually increasing the share of medium-difficulty samples, while decreasing both easy and hard samples, leads to consistent improvements across VideoQA benchmarks. Notably, comparing the second and third rows reveals that while a higher proportion of hard data can benefit reasoning, it may negatively impact general performance. These results highlight the importance of balancing difficulty levels, particularly by increasing medium-difficulty samples, to optimize model training under GRPO-based RLT.

Data Selection for Temporal Grounding. In temporal video grounding, the reward signal is continuous, making intra-group variability a critical factor for effective learning. We use Δ_{IoU} as a proxy for internal diversity within sampled outputs. As shown in Tab. 4b, increasing Δ_{IoU} from 0 to 0.3 consistently improves grounding performance. These results indicate that samples with greater prediction variance provide more informative preference signals, making them more suitable for GRPO-based RLT in temporal grounding tasks.

Generalization of RLT Across Video Tasks. To assess the generalization capability of RLT in video understanding, as shown in 5, we perform both SFT and GRPO-based RLT using only the Charades training set, and evaluate performance on both in-domain and out-of-domain benchmarks. Compared to the SFT baseline, the GRPO-trained model demonstrates substantial performance gains and stronger reasoning capabilities, as evidenced by the generation of explicit "think" traces. Specifically, on the in-domain Charades dataset, our model achieves a +4.7 mIoU improvement. More notably, it generalizes effectively to out-of-domain settings, with a +16.0 mIoU gain on ActivityNet and a +8.3 mIoU gain on the reasoning-focused ActivityNet-RTL benchmark. These results highlight not only the effectiveness of our approach for temporal localization, but also its ability to transfer reasoning capabilities across tasks, thereby enhancing both localization accuracy and temporal reasoning performance. This demonstrates that GRPO-based training equips the model with stronger generalization and adaptability to diverse video domains and task requirements.

Distributional Shifts Before and After RLT. We further examine changes in the model’s output distribution induced by RLT training as shown in 4. For the VideoQA task, we observe that prior to RLT, most samples were classified as easy, though a substantial portion exhibited medium or

Table 4: Ablation Studies: Video QA and TVG Data Selection.

Easy: Middle: Hard	General VideoQA			Reasoning QA	Δ_{IoU}	Charades-STA			
	MVBench	TempCompass	VideoMME	MMVU		Recall@0.3	Recall@0.5	Recall@0.7	mIoU
4 : 4 : 2	64.3	70.0	52.8	59.5	0	78.2	63.9	37.4	54.7
2 : 4 : 4	65.9	70.3	55.9	63.0	0.1	78.0	64.8	38.9	54.9
2 : 6 : 2	67.2	71.3	56.8	62.1	0.2	78.8	63.7	38.5	55.0
1 : 8 : 1	68.1	73.4	57.1	63.4	0.3	78.6	64.5	39.9	55.5
0 : 10 : 0	68.1	72.5	58.6	63.1					

(a) Ablation for Video QA Data Selection

(b) Ablation for TVG Data Selection

Table 5: Temporal Video Grounding OOD Evaluation.

Tuning Type	Charades-STA				ActivityNet				ActivityNet-RTL			
	R@0.3	R@0.5	R@0.7	mIoU	R@0.3	R@0.5	R@0.7	mIoU	R@0.3	R@0.5	R@0.7	mIoU
$\times$	42.4	29.8	14.0	28.0	34.4	22.5	11.6	24.0	7.9	2.6	2.9	6.0
SFT	73.9	61.6	38.5	52.8	33.4	18.9	9.0	23.1	24.0	14.8	7.4	17.8
RLT	80.2	68.3	44.5	57.9	56.9	38.4	20.2	39.1	40.2	22.7	10.9	26.3

hard difficulty. After RLT, the medium-difficulty samples largely shifted toward either easy or hard, indicating a loss in output diversity. For the temporal grounding task, we compare changes in the distribution of the Δ_{IoU} (max–min) and mean IoU scores. Before training, Δ_{IoU} peaked near the center, while mean IoU was skewed leftward. Post-RLT, the Δ_{IoU} peak shifts closer to zero, and the mean IoU distribution moves rightward. This suggests an overall performance improvement, but also a reduction in intra-group variance. These trends reveal important task-specific implications. In the case of VideoQA, the convergence of predictions implies that further RLT provides diminishing returns, as the model’s output becomes overly deterministic and less capable of generating meaningful preference gradients. In contrast, temporal grounding tasks retain sufficient prediction variance post-RLT, indicating continued room for improvement. This highlights the importance of considering output diversity in evaluating the saturation and potential of RLT for different video tasks.

5 Related Work

Video Large Language Models. Video Large Language Models (VideoLLMs) [2, 14, 34, 39, 15] have emerged as unified architectures for video-language understanding. Early works [16, 18] focused on tasks such as video captioning, video question answering, and video-grounded dialogue, typically trained via supervised objectives. More recent efforts [10, 20, 31, 12, 36, 29, 25, 15, 9] have introduced temporally structured datasets and architectures to better capture frame-level alignment and long-range dependencies. Techniques such as sparse sampling and temporal token compression have been proposed to improve temporal perception and efficiency. In parallel, temporal reasoning has gained increasing attention. Lita [11] introduced a benchmark targeting temporal grounding, while Momentor [24] released a large-scale dataset for temporal reasoning. However, most approaches remain grounded in supervised fine-tuning (SFT), with reinforcement learning tuning (RLT) still underexplored for developing temporal understanding and generalizable reasoning in VideoLLMs.

Reinforcement Learning Tuning for MLLMs. Reinforcement learning (RL) has proven effective for improving alignment and reasoning in large language and multimodal models. DeepSeek-R1 [7, 26] showed that preference-based tuning with outcome-level rewards can enable strong chain-of-thought (CoT) reasoning without step-level supervision. Building on this, Kimi-k1.5 [28] adopted rule-based RL strategies to improve reasoning across both text and vision modalities. In the multimodal setting, RLT with verifiable rewards [22, 38, 37, 41] has been used to enhance visual reasoning in tasks such as VQA and image grounding. However, applying RL to video understanding remains limited due to challenges in designing temporally sensitive rewards and modeling spatiotemporal consistency. Only a few works, such as Video-R1 [4], have explored this direction, leaving significant space for further research on reward formulation and training efficiency in the video domain.

6 Conclusion

In this paper, we present an RL-based post-training framework for enhancing the video-specific reasoning abilities of LLMs. We build upon the GRPO algorithm and propose a dual-reward formulation that explicitly supervises both semantic and temporal reasoning. To ensure effective optimization under this framework, we introduce a variance-aware sample selection method based

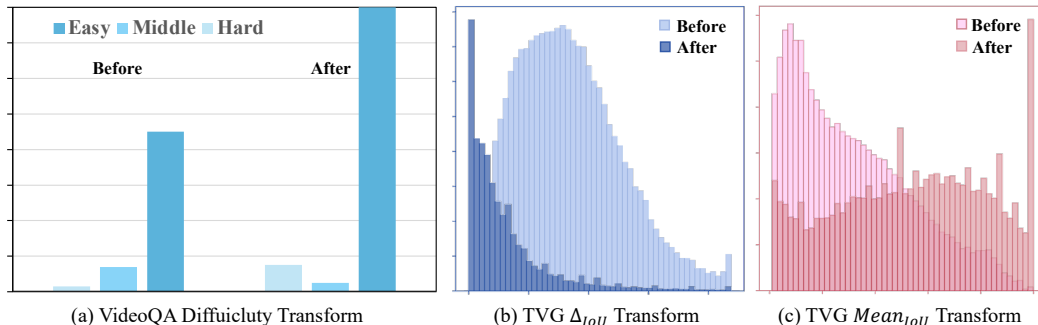


Figure 4: Changes in data distribution before and after RLT. We report the distribution of correctness counts for MC-VideoQA and the IoU variance (max-min) for temporal grounding. The post-RLT data reflects only the subset of samples selected from the original dataset based on our sampling strategy.

on repeated inference and a dynamic sampling strategy that progressively refines the training set. Through the construction of two dedicated datasets and comprehensive experiments on a wide range of video understanding tasks, we demonstrate that our approach significantly improves model performance in reasoning-centric scenarios, even when compared to existing SFT and RLT baselines. Our results highlight the importance of tailored reward design and data selection in post-training VideoLLMs and open up new directions for scalable, reasoning-driven video understanding.

Limitations. This work focuses on reward formulation and data efficiency within the GRPO-based RLT for video understanding. We do not explore the construction of high-quality, fine-grained thinking traces to support SFT for enhancing the model’s explicit thinking and reasoning abilities. Incorporating such structured thinking supervision remains an important direction for future research.

References

- [1] Lisa Anne Hendricks, Oliver Wang, Eli Shechtman, Josef Sivic, Trevor Darrell, and Bryan Russell. Localizing moments in video with natural language. In *Proceedings of the IEEE international conference on computer vision*, pages 5803–5812, 2017.
- [2] Shuai Bai, Keqin Chen, Xuejing Liu, Jialin Wang, Wenbin Ge, Sibao Song, Kai Dang, Peng Wang, Shijie Wang, Jun Tang, et al. Qwen2. 5-vl technical report. *arXiv preprint arXiv:2502.13923*, 2025.
- [3] Zesen Cheng, Sicong Leng, Hang Zhang, Yifei Xin, Xin Li, Guanzheng Chen, Yongxin Zhu, Wenqi Zhang, Ziyang Luo, Deli Zhao, et al. Videollama 2: Advancing spatial-temporal modeling and audio understanding in video-llms. *arXiv preprint arXiv:2406.07476*, 2024.
- [4] Kaituo Feng, Kaixiong Gong, Bohao Li, Zonghao Guo, Yibing Wang, Tianshuo Peng, Benyou Wang, and Xiangyu Yue. Video-r1: Reinforcing video reasoning in mllms. *arXiv preprint arXiv:2503.21776*, 2025.
- [5] Chaoyou Fu, Yuhang Dai, Yondong Luo, Lei Li, Shuhuai Ren, Renrui Zhang, Zihan Wang, Chenyu Zhou, Yunhang Shen, Mengdan Zhang, et al. Video-mme: The first-ever comprehensive evaluation benchmark of multi-modal llms in video analysis. *arXiv preprint arXiv:2405.21075*, 2024.
- [6] Jiyang Gao, Chen Sun, Zhenheng Yang, and Ram Nevatia. Tall: Temporal activity localization via language query. In *Proceedings of the IEEE international conference on computer vision*, pages 5267–5275, 2017.
- [7] Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, et al. Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning. *arXiv preprint arXiv:2501.12948*, 2025.
- [8] Yongxin Guo, Jingyu Liu, Mingda Li, Dingxin Cheng, Xiaoying Tang, Dianbo Sui, Qingbin Liu, Xi Chen, and Kevin Zhao. Vtg-llm: Integrating timestamp knowledge into video llms for enhanced video temporal grounding. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 39, pages 3302–3310, 2025.

- [9] Songhao Han, Wei Huang, Hairong Shi, Le Zhuo, Xiu Su, Shifeng Zhang, Xu Zhou, Xiaojuan Qi, Yue Liao, and Si Liu. Videospresso: A large-scale chain-of-thought dataset for fine-grained video reasoning via core frame selection. *arXiv preprint arXiv:2411.14794*, 2024.
- [10] Bin Huang, Xin Wang, Hong Chen, Zihan Song, and Wenwu Zhu. Vtimellm: Empower llm to grasp video moments. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 14271–14280, 2024.
- [11] De-An Huang, Shijia Liao, Subhashree Radhakrishnan, Hongxu Yin, Pavlo Molchanov, Zhiding Yu, and Jan Kautz. Lita: Language instructed temporal-localization assistant. In *European Conference on Computer Vision*, pages 202–218. Springer, 2024.
- [12] Shaofei Huang, Rui Ling, Hongyu Li, Tianrui Hui, Zongheng Tang, Xiaoming Wei, Jizhong Han, and Si Liu. Unleashing the temporal-spatial reasoning capacity of gpt for training-free audio and language referenced video object segmentation. In *Proceedings of the AAAI Conference on Artificial Intelligence*, 2025.
- [13] Ranjay Krishna, Kenji Hata, Frederic Ren, Li Fei-Fei, and Juan Carlos Niebles. Dense-captioning events in videos. In *Proceedings of the IEEE international conference on computer vision*, pages 706–715, 2017.
- [14] Bo Li, Yuanhan Zhang, Dong Guo, Renrui Zhang, Feng Li, Hao Zhang, Kaichen Zhang, Peiyuan Zhang, Yanwei Li, Ziwei Liu, et al. Llava-onevision: Easy visual task transfer. *arXiv preprint arXiv:2408.03326*, 2024.
- [15] Hongyu Li, Jinyu Chen, Ziyu Wei, Shaofei Huang, Tianrui Hui, Jialin Gao, Xiaoming Wei, and Si Liu. Llava-st: A multimodal large language model for fine-grained spatial-temporal understanding. *arXiv preprint arXiv:2501.08282*, 2025.
- [16] KunChang Li, Yinan He, Yi Wang, Yizhuo Li, Wenhai Wang, Ping Luo, Yali Wang, Limin Wang, and Yu Qiao. Videochat: Chat-centric video understanding. *arXiv preprint arXiv:2305.06355*, 2023.
- [17] Kunchang Li, Yali Wang, Yinan He, Yizhuo Li, Yi Wang, Yi Liu, Zun Wang, Jilan Xu, Guo Chen, Ping Luo, et al. Mvbench: A comprehensive multi-modal video understanding benchmark. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 22195–22206, 2024.
- [18] Yanwei Li, Chengyao Wang, and Jiaya Jia. Llama-vid: An image is worth 2 tokens in large language models. In *European Conference on Computer Vision*, pages 323–340. Springer, 2024.
- [19] Jiajun Liu, Yibing Wang, Hanghang Ma, Xiaoping Wu, Xiaoqi Ma, Xiaoming Wei, Jianbin Jiao, Enhua Wu, and Jie Hu. Kangaroo: A powerful video-language model supporting long-context video input. *arXiv preprint arXiv:2408.15542*, 2024.
- [20] Ruyang Liu, Chen Li, Haoran Tang, Yixiao Ge, Ying Shan, and Ge Li. St-llm: Large language models are effective temporal learners. In *European Conference on Computer Vision*, pages 1–18. Springer, 2024.
- [21] Yuanxin Liu, Shicheng Li, Yi Liu, Yuxiang Wang, Shuhuai Ren, Lei Li, Sishuo Chen, Xu Sun, and Lu Hou. Tempcompass: Do video llms really understand videos? *arXiv preprint arXiv:2403.00476*, 2024.
- [22] Ziyu Liu, Zeyi Sun, Yuhang Zang, Xiaoyi Dong, Yuhang Cao, Haodong Duan, Dahua Lin, and Jiaqi Wang. Visual-rft: Visual reinforcement fine-tuning. *arXiv preprint arXiv:2503.01785*, 2025.
- [23] Viorica Patraucean, Lucas Smaira, Ankush Gupta, Adria Recasens, Larisa Markeeva, Dylan Banarse, Skanda Koppula, Mateusz Malinowski, Yi Yang, Carl Doersch, et al. Perception test: A diagnostic benchmark for multimodal video models. *Advances in Neural Information Processing Systems*, 36:42748–42761, 2023.

- [24] Long Qian, Juncheng Li, Yu Wu, Yaobo Ye, Hao Fei, Tat-Seng Chua, Yueting Zhuang, and Siliang Tang. Momentor: Advancing video large language model with fine-grained temporal reasoning. *arXiv preprint arXiv:2402.11435*, 2024.
- [25] Shuhuai Ren, Linli Yao, Shicheng Li, Xu Sun, and Lu Hou. Timechat: A time-sensitive multimodal large language model for long video understanding. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 14313–14323, 2024.
- [26] Zhihong Shao, Peiyi Wang, Qihao Zhu, Runxin Xu, Junxiao Song, Xiao Bi, Haowei Zhang, Mingchuan Zhang, YK Li, Y Wu, et al. Deepseekmath: Pushing the limits of mathematical reasoning in open language models. *arXiv preprint arXiv:2402.03300*, 2024.
- [27] Haozhan Shen, Peng Liu, Jingcheng Li, Chunxin Fang, Yibo Ma, Jiajia Liao, Qiaoli Shen, Zilun Zhang, Kangjia Zhao, Qianqian Zhang, et al. Vlm-r1: A stable and generalizable r1-style large vision-language model. *arXiv preprint arXiv:2504.07615*, 2025.
- [28] Kimi Team, Angang Du, Bofei Gao, Bowei Xing, Changjiu Jiang, Cheng Chen, Cheng Li, Chenjun Xiao, Chenzhuang Du, Chonghua Liao, et al. Kimi k1. 5: Scaling reinforcement learning with llms. *arXiv preprint arXiv:2501.12599*, 2025.
- [29] Haibo Wang, Zhiyang Xu, Yu Cheng, Shizhe Diao, Yufan Zhou, Yixin Cao, Qifan Wang, Weifeng Ge, and Lifu Huang. Grounded-videollm: Sharpening fine-grained temporal grounding in video large language models. *arXiv preprint arXiv:2410.03290*, 2024.
- [30] Yi Wang, Xinhao Li, Ziang Yan, Yinan He, Jiashuo Yu, Xiangyu Zeng, Chenting Wang, Changlian Ma, Haian Huang, Jianfei Gao, et al. Internvideo2. 5: Empowering video mllms with long and rich context modeling. *arXiv preprint arXiv:2501.12386*, 2025.
- [31] Yueqian Wang, Xiaojun Meng, Jianxin Liang, Yuxuan Wang, Qun Liu, and Dongyan Zhao. Hawkeye: Training video-text llms for grounding text in videos. *arXiv preprint arXiv:2403.10228*, 2024.
- [32] Junbin Xiao, Xindi Shang, Angela Yao, and Tat-Seng Chua. Next-qa: Next phase of question-answering to explaining temporal actions. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 9777–9786, 2021.
- [33] Junbin Xiao, Angela Yao, Yicong Li, and Tat-Seng Chua. Can i trust your answer? visually grounded video question answering. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 13204–13214, 2024.
- [34] Jiabo Ye, Haiyang Xu, Haowei Liu, Anwen Hu, Ming Yan, Qi Qian, Ji Zhang, Fei Huang, and Jingren Zhou. mplug-owl3: Towards long image-sequence understanding in multi-modal large language models. *arXiv preprint arXiv:2408.04840*, 2024.
- [35] En Yu, Kangheng Lin, Liang Zhao, Yana Wei, Zining Zhu, Haoran Wei, Jianjian Sun, Zheng Ge, Xiangyu Zhang, Jingyu Wang, et al. Unhackable temporal rewarding for scalable video mllms. *arXiv preprint arXiv:2502.12081*, 2025.
- [36] Xiangyu Zeng, Kunchang Li, Chenting Wang, Xinhao Li, Tianxiang Jiang, Ziang Yan, Songze Li, Yansong Shi, Zhengrong Yue, Yi Wang, et al. Timesuite: Improving mllms for long video understanding via grounded tuning. *arXiv preprint arXiv:2410.19702*, 2024.
- [37] Yufei Zhan, Yousong Zhu, Shurong Zheng, Hongyin Zhao, Fan Yang, Ming Tang, and Jinqiao Wang. Vision-r1: Evolving human-free alignment in large vision-language models via vision-guided reinforcement learning. *arXiv preprint arXiv:2503.18013*, 2025.
- [38] Jingyi Zhang, Jiaxing Huang, Huanjin Yao, Shunyu Liu, Xikun Zhang, Shijian Lu, and Dacheng Tao. R1-vl: Learning to reason with multimodal large language models via step-wise group relative policy optimization. *arXiv preprint arXiv:2503.12937*, 2025.
- [39] Peiyuan Zhang, Kaichen Zhang, Bo Li, Guangtao Zeng, Jingkang Yang, Yuanhan Zhang, Ziyue Wang, Haoran Tan, Chunyuan Li, and Ziwei Liu. Long context transfer from language to vision. *arXiv preprint arXiv:2406.16852*, 2024.

- [40] Yilun Zhao, Lujing Xie, Haowei Zhang, Guo Gan, Yitao Long, Zhiyuan Hu, Tongyan Hu, Weiyuan Chen, Chuhan Li, Junyang Song, et al. Mmvu: Measuring expert-level multi-discipline video understanding. *arXiv preprint arXiv:2501.12380*, 2025.
- [41] Hengguang Zhou, Xirui Li, Ruochen Wang, Minhao Cheng, Tianyi Zhou, and Cho-Jui Hsieh. R1-zero's "aha moment" in visual reasoning on a 2b non-sft model. *arXiv preprint arXiv:2503.05132*, 2025.

Table 6: Result for Treating VideoQA as Grounded QA.

Infer as GQA	General VideoQA			Reasoning QA
	MVBench	TempCompass	VideoMME	MMVU
$\times$	68.1	73.3	57.6	65.0
$\checkmark$	67.8	72.8	57.3	63.5

A Data Selection Implementation Details

To filter high-quality data suitable for GRPO training from the Temporal-RLT-Full-490k dataset, we first categorize the data into three types: tasks requiring **Discrete Rewards**, tasks requiring **Continuous Rewards**, and tasks requiring both. These categories align with grounded VideoQA tasks. Initially, we utilize the entire dataset to perform "**No-Think**" inference using the Qwen-VL-2.5-7B model to generate preliminary results. Subsequently, we apply task-specific processing for each category, using the same settings as the GRPO inference process (temperature = 1.0, top-p = 0.99, and sample-N = 8).

Discrete Rewards Data Selection For discrete rewards data, we set thresholds of $\tau_{\text{easy}} = 1$ and $\tau_{\text{hard}} = 7$ to classify the data into easy, medium, and hard categories. Based on our ablation study, which showed that medium-level data performs best for multi-choice tasks, we select 16k samples from the medium category. Additionally, we ensure an even distribution of samples from each data source within the dataset.

Continuous Rewards Data Selection For continuous rewards data, our ablation study revealed that a larger Δ_{IoU} leads to better results. Therefore, we select data with $\Delta_{\text{IoU}} \geq 0.3$. From this subset, we extract 8k samples while maintaining an even distribution of samples across all data sources.

Grounded QA Data Selection For tasks requiring both types of rewards, we apply the aforementioned rules simultaneously. However, due to the limited availability of data in this category, we lower the Δ_{IoU} threshold to 0.1. From this filtered subset, we select 8k samples, ensuring an even distribution across all data sources.

By following these steps, we ensure the selection of high-quality, balanced data tailored for GRPO training, optimizing performance across various task types. As a result, we constructed the **Temporal-RLT-32k** dataset, which serves as a refined subset designed to enhance the training process and improve task-specific outcomes.

B Results

B.1 Treating VideoQA as Grounded QA

Our model inherently supports treating VideoQA as a form of Grounded QA reasoning, which can be achieved simply by adjusting the prompt format. We evaluated the approach of treating VideoQA as Grounded QA, and while there is some performance degradation as shown in 6, this method allows the model to specify highly relevant time intervals being observed before answering the question.

B.2 Visualization

We visualized the example output results on the tasks of Temporal Video Grounding 5, Video QA 6, and Grounded QA 7, including both the reasoning process and the final answers.

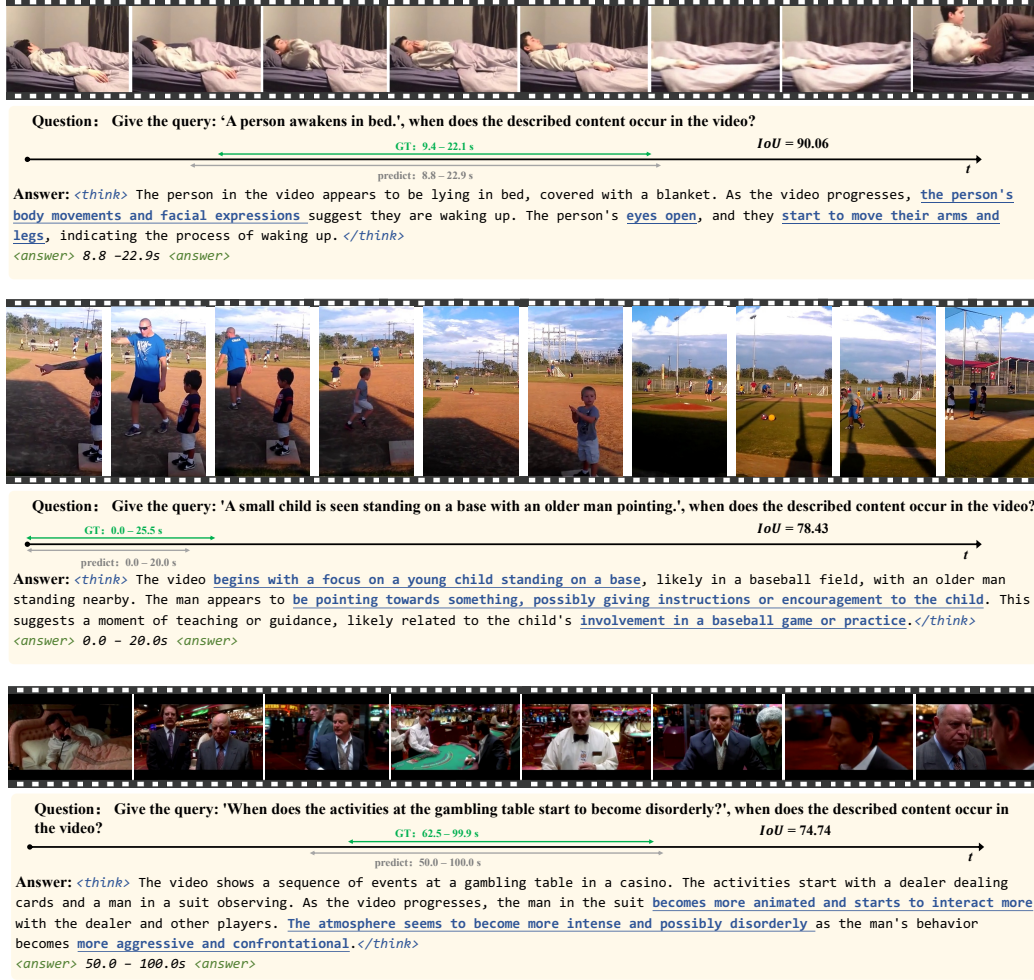
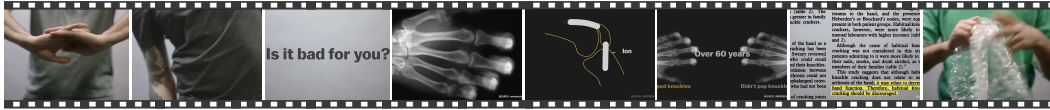


Figure 5: Visualization of Temporal-RLT’s reasoning outputs on the Temporal Video Grounding task. The three examples, from top to bottom, come from Charades-STA, ActivityNet, and ActivityNet-RTL, respectively.



Question: Which of the following could be the main topic of the video segment?

- ☒ (A) The surgical procedure for joint replacement.
- ☒ (B) The nutritional requirements for healthy bones.
- ☒ (C) The explanation of how joint lubrication works.
- ☒ (D) What happens to your knuckles when you crack them.

Answer: *<think>* The video segment begins with a question about what happens to our joints when we crack them and whether it is bad for you. It then delves into the explanation of how joint lubrication works, using a diagram of a synovial joint and a skeleton to illustrate the process. The video also touches on the effects of habitual knuckle cracking on hand function and the results of a study over 60 years. The segment concludes with a call to action to subscribe to the channel and watch more videos. *</think>*

<answer> (D) *<answer>*



Question: How many times did the person show objects to the camera?

- ☒ (A) 5
- ☒ (B) 6
- ☒ (C) 8

Answer: *<think>* The person shows the following objects to the camera:

1. A clear glass with a blue label.
2. A green glass.
3. A red wallet.
4. A red book.
5. A black and white book.
6. A red book again. *</think>*

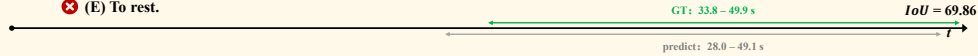
<answer> (B) *<answer>*

Figure 6: Visualization examples of Temporal-RLT’s reasoning output on Video QA task. The two examples are from VideoMME and MVBench, respectively.



Question: Why does the brown cat watch the other cat eat food?

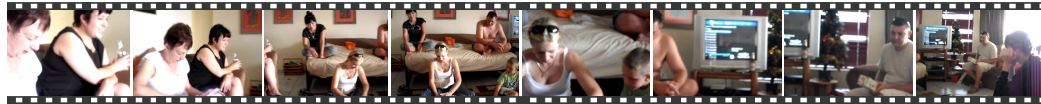
- ☒ (A) Wants to go into box.
- ☒ (B) Resting.
- ☒ (C) Change playing position.
- ☒ (D) Waiting for his turn.
- ☒ (E) To rest.



Answer: *<think>* The brown cat is watching the other cat eat food, which suggests that the brown cat might be waiting for its turn to eat or might be observing the other cat’s behavior. There is no indication that the brown cat wants to go into a box or rest, and the other cat’s actions do not suggest a change in playing position or a desire to rest. *</think>*

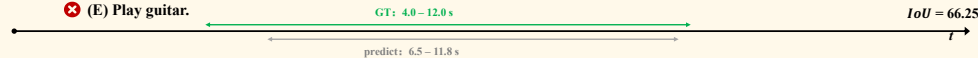
<observe> 28.0 - 49.1s *</observe>*

<answer> (D) *<answer>*



Question: What does the lady in white on the floor do as everyone was sitting around?

- ☒ (A) Go to lady in stripes.
- ☒ (B) Unwrap something.
- ☒ (C) Adjust the girl’s jacket.
- ☒ (D) Shuffle cards.
- ☒ (E) Play guitar.



Answer: *<think>* The lady in white on the floor appears to be adjusting or interacting with something on the floor, possibly a piece of clothing or an object, as she is sitting on the floor and leaning forward. The other individuals in the room are mostly sitting and appear to be watching or engaged in their own activities, but there is no clear indication of them playing a guitar or shuffling cards. The lady in stripes is not visible in the video, and the girl’s jacket is not mentioned as a focus of the video. *</think>*

<observe> 6.5 - 11.8s *</observe>*

<answer> (B) *<answer>*

Figure 7: Visualization of Temporal-RLT’s reasoning output on Grounded QA task. The examples are from the test set of NextGQA.