

Question-Instructed Visual Descriptions for Zero-Shot Video Question Answering

Anonymous ACL submission

Abstract

We present Q-ViD, a simple approach for video question answering (video QA), that unlike prior methods, which are based on complex architectures, computationally expensive pipelines or use closed models like GPTs, Q-ViD relies on a single instruction-aware open vision-language model (InstructBLIP) to tackle videoQA using frame descriptions. Specifically, we create captioning instruction prompts that rely on the target questions about the videos and leverage InstructBLIP to obtain video frame captions that are useful to the task at hand. Subsequently, we form descriptions of the whole video using the question-dependent frame captions, and feed that information, along with a question-answering prompt, to a large language model (LLM). The LLM is our reasoning module, and performs the final step of multiple-choice QA. Our simple Q-ViD framework achieves competitive or even higher performances than current state of the art models on a diverse range of videoQA benchmarks, including NExT-QA, STAR, How2QA, TVQA and IntentQA. Our code will be publicly available at:

1 Introduction

Recently, Vision-Language models have shown remarkable performances in image question-answering tasks (Goyal et al., 2017; Marino et al., 2019; Schwenk et al., 2022), with models such as Flamingo (Alayrac et al., 2022), BLIP-2 (Li et al., 2023b), InstructBlip (Dai et al., 2023) and mPLUG-Owl (Ye et al., 2023) showing strong reasoning capabilities in the vision-language space. Image captioning (Vinyals et al., 2015; Ghandi et al., 2023) is one of the capabilities in which these models truly excel, as they can generate detailed linguistic descriptions from images. Different works have leveraged this capability in many ways for zero-shot image-question answering, such as giving linguistic context to images (Hu et al.,

2022; Ghosal et al., 2023), addressing underspecification problems in questions (Prasad et al., 2023), coordination of multiple image captions to complement information (Chen et al., 2023b), or by combining captions with other type of linguistic information from the image (Berrios et al., 2023). In this manner, the reasoning capabilities of LLMs can be directly used to reason about the linguistic image descriptions and generate an answer for the given visual question.

This approach has been successful for images, but in the case of video-question answering tasks (Lei et al., 2018; Li et al., 2020b; Xiao et al., 2021; Wu et al., 2021; Li et al., 2023a) this is more challenging. Video possesses multiple image frames that have relationships between each other and involve the recognition of objects, actions, as well as the inference about semantic, temporal, causal reasoning and much more (Zhong et al., 2022). Thus, some works (Chen et al., 2023a; Wang et al., 2023) have focused on using powerful LLMs like ChatGPT to either ask visual questions to image-language models like BLIP-2 or to respond and retrieve useful information from large datasets with detailed information from the video. Similarly, Zhang et al., (2023a) have leveraged the reasoning capabilities of GPT-3.5 to create textual summaries from the video, and later perform video QA using only textual information. While others (Wang et al., 2022b; Zeng et al., 2022) combine linguistic information from multiple sources such as captions, visual tokenization or even subtitles of input speech. In summary, current methods for video QA rely on any combination of closed LLMs, expensive training regimes, and complex architectures with multiple modules (Yang et al., 2022; Ko et al., 2023; Yu et al., 2023; Momeni et al., 2023; Li et al., 2023c; Zhang et al., 2023a). In contrast, we introduce Q-ViD a simple Question-Instructed Visual Descriptions for video QA approach that relies on an instruction-aware vision-language model,

InstructBLIP (Dai et al., 2023), to automatically generate rich captions from the frames. In this manner, we effectively turn the vQA task into a text QA task. More specifically, given an input video V we sample n number of frames, then, we generate question-specific instructions to prompt the multimodal instruction tuned model to generate captions for each frame. Afterwards, we form a video description by concatenating all the generated question-dependent captions from InstructBLIP, and use it along with the question, options and an instruction prompt as input to the LLM-based reasoning module that generates an answer to the multiple-choice question about the video. We demonstrate the effectiveness of Q-ViD on five challenging multiple choice video question answering tasks (NExT-QA, STAR, How2QA, TVQA, IntentQA), showing that this simple framework can achieve strong performances comparable with more complex pipelines. Our contributions are summarized as follows:

- We propose Q-ViD, a simple gradient-free approach for zero-shot video QA that relies on an open instruction-tuned multimodal model to extract question-specific descriptions of frames to transform the video QA task into a text QA one.
- Our approach achieves strong zero-shot performance that is competitive or even superior to more complex architectures such as SeViLa, Internvideo, and Flamingo. It even compares favorably with recent solutions that include GPT APIs, like LloVi and ViperGPT.

2 Related Work

2.1 Multimodal Pretraining for Video QA

The strong reasoning capabilities of LLMs (Chung et al., 2022; Touvron et al., 2023; Brown et al., 2020; Hoffmann et al., 2022) in natural language processing tasks has motivated to apply these models for visual understanding. Currently, LLMs have been successfully adapted to understand images (Li et al., 2023b; Ye et al., 2023; Chen et al., 2023c), but applying the same principles for video is more challenging. Approaches for VideoQA rely on image-language models, and adapt those to process video by using fixed amounts of video frames as input (Alayrac et al., 2022; Yu et al., 2023; Yang et al., 2022), or by selecting key-frames from the initial sequence (Yu et al., 2023; Li et al., 2023c).

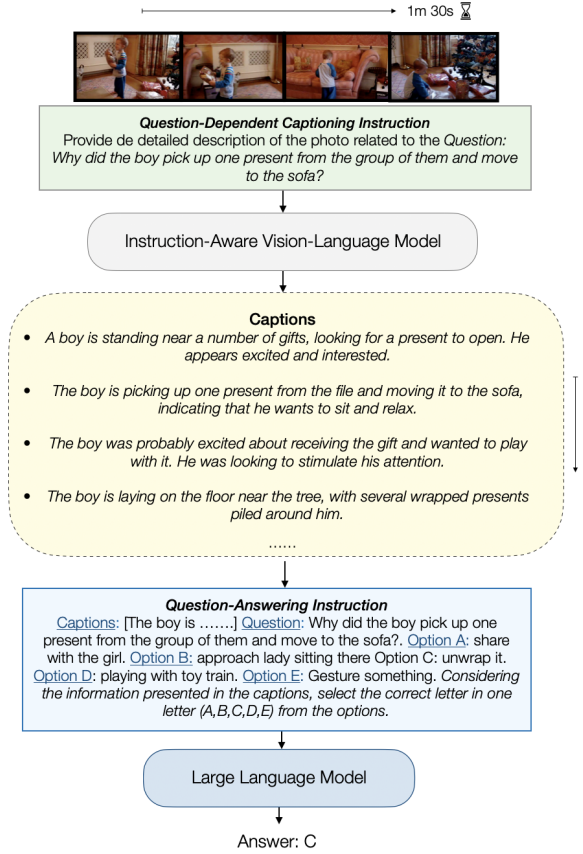


Figure 1: **Overview of Q-ViD.** We propose relying on a instructed-tuned multimodal model to generate question-dependent frame captions to perform video QA using text. This simple approach achieves competitive results with more complex architectures or GPT-based methods

Commonly, these works use frozen visual and language models and focus only on modality alignment. Models like Flamingo (Alayrac et al., 2022) uses a fixed amount of video frames as input and bridges modalities by training a perceiver resampler and gated attention layers between Chinchilla LLMs (Hoffmann et al., 2022). While others, like SeViLa (Yu et al., 2023) relies on BLIP-2 (Li et al., 2023b) for modality alignment, using an intermediate pretrained module called Q-former. SeViLa, first perform key-frame localization and then video QA with Flan-T5 LLMs (Chung et al., 2022). On the other hand, other works apart from using frozen vision models, adapt the LLM to visual inputs using adapter tokens (Zhang et al., 2023b) or intermediate trainable modules (Houlsby et al., 2019). Models like Flipped-VQA (Ko et al., 2023) focuses on adapting LLaMa (Touvron et al., 2023) to video QA by using adapter tokens along with different training objectives to leverage the temporal and causal reasoning abilities of LLMs. Similarly, Frozen-

Bilm (Yang et al., 2022) exploit the strong zero-shot performance of BILM, a frozen bidirectional language model that is adapted to video QA by using lightweight trainable modules. Despite the success of all these models, they require complex architectures and training regimes, unlike these works we build a simple, gradient-free, approach for zero-shot video QA.

2.2 Image Captions for Video Understanding

One of the core strengths of image-language models (Alayrac et al., 2022; Li et al., 2023b; Dai et al., 2023) is the generation of image captions, thus due to the current strong zero-shot capabilities of LLMs, captions can be directly use to reason about visual content. This has been successfully leveraged in the image-language space for image question-answering with approaches such as Lens (Berrios et al., 2023), Img2LLM (Guo et al., 2023) and PromptCat (Hu et al., 2022) that gather image captions and other type of linguistic information to answer a visual question. While similar approaches have been taken for videos, the use of large models like GPTs is very common, with models such as ChatCaptioner (Chen et al., 2023a), ViperGPT (Surís et al., 2023), ChatVideo (Wang et al., 2023), VidIL (Wang et al., 2022b), Socratic Models (Zeng et al., 2022), and LLoVi (Zhang et al., 2023a) have been applied for video-language tasks, common methods use GPTs to either interact with image-language models to get visual descriptions, or to make summaries from captions and other type of information such as visual tokenization, subtitles of speech and more. Unlike these approaches, we do not use GPTs or multiple computationally expensive modules in any part of our pipeline to achieve strong zero-shot performance on video QA.

3 Method

Recently, vision-language models trained with instruction tuning (Dai et al., 2023; Zhu et al., 2023; Liu et al., 2023) have shown impressive capabilities to faithfully follow instructions and extract visual representations adapted to the task at hand. Thus, with Q-ViD (Fig. 2), we propose to leverage these capabilities for multiple-choice video QA, and turn this task into textual QA using InstructBLIP (Dai et al., 2023). We use a question-dependent captioning prompt as the input instruction, to guide InstructBLIP to generate video frame descriptions that are more relevant for the given question. Af-

terwards, we reuse the LLM from InstructBLIP and use it as our reasoning module. This LLM (Flan-T5) takes a question-answering prompt as input, that consists of a video description formed by the concatenation of all the question-dependent frame captions, the question, options and a task instruction. Considering that Flan-T5 is also originally trained with instructions, we aim to leverage its reasoning capabilities to correctly answer the question given only the text we just described as input. Our simple approach does not rely on complex pipelines or closed GPT models, which makes it easy, cheaper and straight forward to use for zero-shot video QA. On the other hand, Q-ViD is flexible and model agnostic, which means we can use any multimodal models available. This section presents our approach in detail. First, we introduce some preliminary information on InstructBLIP, which serves as the foundation of our work, and then we provide a detailed overview for all components from our Q-ViD framework.

3.1 Preliminaries: InstructBLIP

We rely on InstructBLIP (Dai et al., 2023) as the foundational architecture of Q-ViD. InstructBLIP is a vision-language instruction tuning framework based on a Query Transformer (Q-former) and frozen vision and language models. Unlike BLIP-2 (Li et al., 2023b), which is based on an instruction-agnostic approach, InstructBLIP can obtain visual features depending on specific instructions of the task at hand using an instruction-aware Q-former, which in addition to query embeddings, uses instruction tokens to guide the Q-former in extracting specific image features. Subsequently, a LLM (Flan-T5) uses these features to generate visual descriptions depending on the input instructions. In our approach we adopt this model to obtain frame captions that are dependant on the video questions, thus, we aim to gather the most important information from each part of the video and use it as input for our reasoning module to answer the given question. Because of our Q-ViD framework is a zero-shot approach, we do not train any part of InstructBLIP, and keep all of its parts frozen.

3.2 Q-ViD: Generating Frame Descriptions for Video QA

We focus on automatically generating meaningful captions that can provide enough information about what is happening in the video to the LLM. We assume that if captions for the frames con-

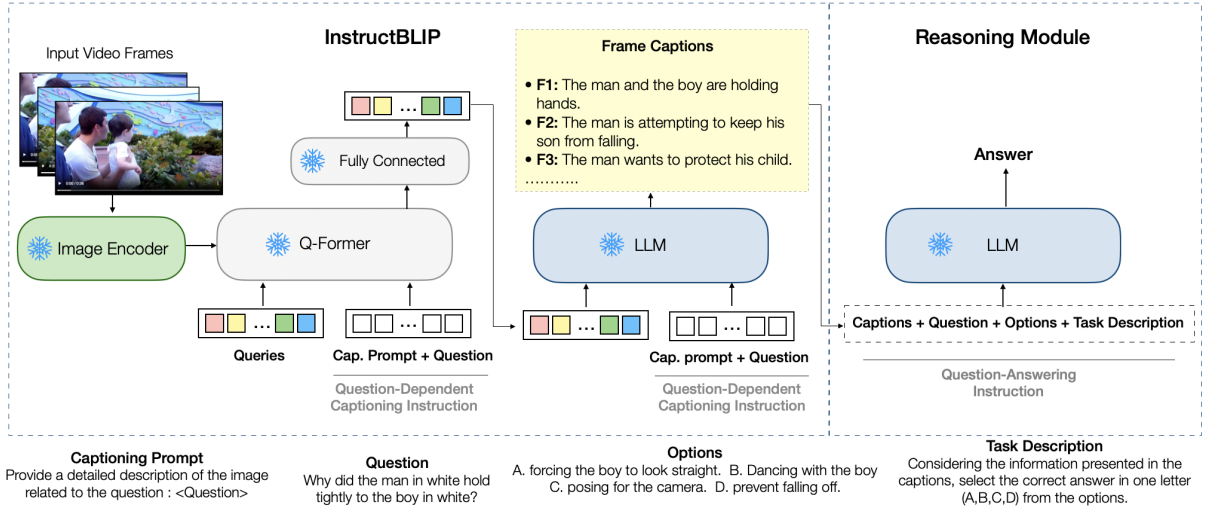


Figure 2: **Our pipeline for Zero-shot Video QA.** Q-ViD prompts InstructBLIP, to obtain video frame descriptions that are tailored to the question needing answer.

tain relevant information related to the question needing answer, then an LLM should be able to answer the question correctly without additional need for frame/video input. As shown in Fig. 2, given an input video V , we use a uniform sampling strategy and extract a set of n video frames $\{f_1, f_2, \dots, f_n\}$. We then use InstructBLIP, referred as I_b , to obtain instruction-aware visual captions c_i for each frame f_i , as follows $c_i = I_b(f_i, E)$, where E represents the question-dependent captioning instruction. Q-ViD automatically generates E by concatenating a request for a caption, referred as B (e.g. "Provide a detailed description of the image related to the question:") and a question, referred as Q (e.g. "Why did the man in white held tightly to the boy in white?"), represented as follows $E = \text{concat}(B, Q)$. Specifically, E is used as input to Q-former and to the LLM of InstructBLIP to obtain specific visual representations and frame descriptions respectively. Thus, we represent the input video V as a set of question-dependent frame captions $c = (c_1, c_2, \dots, c_n)$, where each caption is conformed by a sequence of w_m words $c_i = (w_1, w_2, \dots, w_m)$. In this way, we extract textual information from V , the captions, that is going to be useful for the question answering task. Next, we describe the reasoning module of Q-ViD and how these question-dependent captions are used to perform video QA.

3.3 Q-ViD: Reasoning Module

We reuse the frozen LLM (Flan-T5) from InstructBLIP and implement it as the reasoning module of

Q-ViD. In order to perform video QA using language, we first concatenate the question-dependent frame captions $C = [c_1, c_2, \dots, c_n]$ in the same order they appear in the video. Then, we create a question-answering instruction L as follows: $L = \text{concat}(C, Q, A, T)$. In other words, we concatenate in L the list of captions C , question Q , possible answers A and a task description T (e.g. "Considering the information presented in the captions, select the correct answer in one letter (A,B,C) from the options."). Our goal is to leverage the LLM reasoning linguistic capabilities by providing a set of captions that were tailored to be relevant for the specific question Q . Our experiments in Section 4, show that this simple approach works surprisingly well, showing to be competitive, and even superior in some cases, in comparison with more complex pipelines. Next, we describe in more detail the prompts used for question-dependent captioning and video QA.

3.4 Q-ViD: Prompt Design

First, to get question-dependent captions for each frame, given the question Q we prompt InstructBLIP with a question-dependent captioning instruction: "Provide a detailed description of the image related to the question: $\{Q\}$ ". This instruction is used along with queries as input to the frozen Q-Former and LLM modules of InstructBLIP to extract specific visual features and generate question-dependent descriptions. Afterwards, to perform QA with the reasoning module, given the list of captions C and the list of possible an-

swers $A = [a_1, \dots, a_m]$ with m being the number of options provided in each dataset, we prompt the language model as follows: "*Captions: {C} Question: {Q}. Option A: a_1 Option B: a_2 Option C: a_3 . Considering the information presented in the captions, select the correct answer in one letter from the options (A,B,C)*". In this prompt, in addition to C , Q and A , we added a small instruction at the end to specify in detail that a single letter is needed as output.

4 Experiments

In this section, we present our experiments for zero-shot video QA. First, we describe the datasets we used and the implementation details. Then, we evaluate our approach, compare Q-ViD with other state of the art models for video QA and provide a comprehensive analysis of the model's performance. Lastly, we conduct some ablation studies of Q-ViD regarding the instructions prompt design.

4.1 Datasets

To test our approach, we conduct experiments on the following multiple-choice video QA benchmarks. To make comparisons with prior work we use the validation set in NExT-QA, STAR, How2QA and TVQA, meanwhile in IntentQA we use the test set. More details are shown below:

- **NExT-QA** (Xiao et al., 2021): A benchmark focused on Temporal, Causal and Descriptive reasoning type of questions. Contains 5,440 videos and 48K multiple-choice questions in total. We perform our experiments using the validation set that is conformed by 570 videos and 5K multi-choice questions.
- **STAR** (Wu et al., 2021): A benchmark that evaluates situated reasoning in real-world videos, is focused on interaction, sequence, prediction and feasibility type of questions. It contains 22K situation video clips and 60K questions. We perform evaluations on the validation set with 7K multiple-choice questions.
- **HOW2QA** (Li et al., 2020a): A dataset that consists on 44K question-answering pairs for 22 thousand 60-second clips selected from 9035 videos. We perform experiments on the validation set with 2.8K questions.
- **TVQA** (Lei et al., 2018): A large scale video QA dataset based on six popular TV shows. It has 152K multiple-choice questions and 21K video clips. For our zero-shot evaluations we

use the validation set with 15K video-question pairs.

- **IntentQA** (Li et al., 2023a): A dataset focused on video intent reasoning. It contains 4K videos and 16K multiple-choice question-answer samples. In this case, we use the test set for our zero-shot evaluations which contains 2K video-question answering samples.

4.2 Implementation Details

For Q-ViD we adopt InstructBLIP-Flan-T5_{XXL} with 12.1B parameters, as a default vision encoder it uses ViT-g/14 (Fang et al., 2023), and as language model FlanT5_{XXL} (Chung et al., 2022). We extract 64 frames per video, as in preliminary experiments this number worked well. For frame captioning, we use a maximum number of 30 tokens per description and top-p sampling with $top_p = 0.7$ to get varied captions. Regarding our reasoning module, we reuse and adopt the corresponding Flan-T5 language model from InstructBLIP. In this case we do not use top-p sampling. Our experiments were conducted using 4 NVIDIA A100 (40GB) GPUs using the Language-Vision Intelligence library LAVIS (Li et al., 2022) and the released code from SeViLa (Yu et al., 2023).

4.3 Overall Performance

Table 1 provides a detailed overview on the performance of Q-ViD on the validation set of NExT-QA, STAR, HOW2QA and TVQA. We compare our approach with current state of the art methods such as SeViLa (Yu et al., 2023), FrozenBILM (Yang et al., 2022) and VideoChat2 (Li et al., 2024), as well as, with GPT-based models like ViperGPT (Surís et al., 2023) and LLoVi (Zhang et al., 2023a). The results obtained from our experiments demonstrate the surprisingly competitive nature of Q-ViD, outperforming or being competitive with previous methods with more complex architectural pipelines such as SeViLa, VideoChat2 and LLoVi. For fair comparisons, we gray out methods that use GPTs.

Specifically, on **NExT-QA**, Q-ViD outperforms SeViLa by 2.7% of average accuracy, and achieves almost the same state of the art results of LLoVi, a framework based of GPT-3.5. Notably, Q-ViD is the best-performing model on causal questions, temporal questions, and overall average performance among methods that are not based on GPTs, showing the ability of this approach to perform action reasoning, which is the target of NExT-QA. With **STAR**, Q-ViD achieves the second

Models	NExT-QA				STAR					How2QA	TVQA
	Tem.	Cau.	Des.	Avg.	Int.	Seq.	Pre.	Fea.	Avg.		
<i>GPT-Based Models</i>											
ViperGPT (Surís et al., 2023)	-	-	-	60.0	-	-	-	-	-		
LLoVi (Zhang et al., 2023a)	61.0	69.5	75.6	67.7	-	-	-	-	-		
Flamingo-9B (Alayrac et al., 2022)	-	-	-	-	-	-	-	-	41.8		
Flamingo-80B (Alayrac et al., 2022)	-	-	-	-	-	-	-	-	39.7		
FrozenBILM (Yang et al., 2022)	-	-	-	-	-	-	-	-	-	41.9	29.7
VFC (Momeni et al., 2023)	51.6	45.4	64.1	51.6	-	-	-	-	-		
InternVideo (Wang et al., 2022a)	48.0	43.4	65.1	49.1	43.8	43.2	42.3	37.4	41.6	62.2	35.9
BLIP-2 ^{voting} (Yu et al., 2023)	59.1	61.3	74.9	62.7	41.8	39.7	40.2	39.5	40.3	69.8	35.7
BLIP-2 ^{concat} (Yu et al., 2023)	59.7	60.8	73.8	62.4	45.4	41.8	41.8	40.0	42.2	70.8	36.6
SeViLa (Yu et al., 2023)	61.3	61.5	75.6	63.6	48.3	45.0	44.4	40.8	44.6	72.3	38.2
VideoChat2 (Li et al., 2024)	57.4	61.9	69.9	61.7	58.4	60.9	55.3	53.1	59.0	-	40.6
Q-ViD (Ours)	61.6	67.6	72.2	66.3	48.2	47.2	43.9	43.4	45.7	71.4	41.0

Table 1: **Zero-shot results on video question answering.** For fair comparison we gray out methods that rely on closed GPTs. We bold the best results, and underline the second-best results. QViD shows to be competitive and even outperform some more complex frameworks for zero-shot video QA.

best average accuracy behind VideoChat2, outperforming all other methods like SeViLa by **1.1%**, BLIP-2^{concat} by **3.5%**, InternVideo by **4.1%** and Flamingo-80B by **6%**. Also note that Q-ViD achieves the second best performances on sequence and feasibility type of question of STAR. Lastly, on **How2QA** we achieve the second best performance behind SeViLa, and achieves the best overall performance for **TVQA** with an improvement of **0.4%** to the previous best-performing method VideoChat2.

On the other hand, in Table 2 we evaluate our approach on **IntentQA**, we use the test set of this benchmark in order to compare with prior works. We take the same comparison made from (Zhang et al., 2023a), and divide the table in two categories, Supervised and Zero-shot approaches. Q-ViD continues showing strong results, greatly outperforming all supervised methods and the SeViLa zero-shot performance by **2.7%**. Interestingly, Q-ViD almost achieves the best overall performance from the GPT-based method Llovi. These results demonstrate that our approach can be used among different video QA tasks and be able to achieve strong zero-shot performances.

4.4 Ablation Studies

In this section, we perform some ablation studies related to the instruction prompt selection for Q-ViD. For these experiments, we chose NExT-QA and STAR as our benchmarks, and report results on the validation sets on each dataset. Specifically, we test two model variations, using InstructBLIP-FlanT5_{XL} (Q-ViD_{XL}) and the one used to report our main results, using InstructBLIP-FlanT5_{XXL} (Q-ViD_{XXL}), we test different prompts to analyze

Models	Acc.(%)
<i>Supervised</i>	
HQGA (Surís et al., 2023)	47.7
VGT (Alayrac et al., 2022)	51.3
BlindGPT (Alayrac et al., 2022)	51.6
CaVIR (Alayrac et al., 2022)	57.6
<i>Zero-shot</i>	
SeViLa (Yu et al., 2023)	<u>60.9</u>
LLoVi (Zhang et al., 2023a)	64.0
Q-ViD (Ours)	63.6

Table 2: **Performance on IntentQA.** Q-ViD shows to outperform supervised approaches, strong zero-shot baselines like SeViLa and obtain almost the same performance from the GPT-based model LLoVi.

and compare the use of question-dependent and general descriptive captions. Additionally, we also make some ablation experiments for the question-answering instruction prompt that is used by the reasoning module to perform multi-choice QA. We discuss our findings in detail below.

4.4.1 Prompt Analysis

We focus on analyzing the impact on performance of the Captioning and QA templates of Q-ViD. First, for captioning templates (Fig.3) we compare two variants: (1) General prompts and (2) Question-dependent prompts. With general prompts we focus on obtaining general descriptions of the images (frames) and with question-dependent prompts on information related to the question of the task at hand. In order to leverage as much as possible the instruction-based capabilities of InstructBLIP we create these prompts based on similar templates to

the ones used by InstructBLIP in its original training setup. For these experiments we use the Base Question-Answering prompt that is used as input of the reasoning module to perform multi-choice QA.

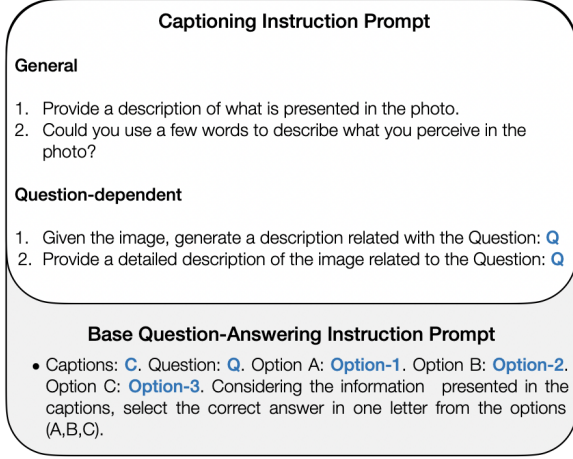


Figure 3: **Variation of captioning templates.** We focus on comparing general and question-dependent captioning prompts (Top). For both cases we use the same Base QA instruction prompt (Bottom).

Table 3 compares the performance of Q-ViD_{XL} and Q-ViD_{XXL} using the general, and question-dependent captioning prompts. It can be seen that performance varies between both models, Q-ViD_{XL} achieves better performances by using general descriptive prompts. When comparing its best variants using the (2)General and (1)Dependent templates, the former further increases the average accuracy by +1.4% on NExT-QA and +3.1% on STAR. On the other hand, the same behaviour is not shown using a bigger model, Q-ViD_{XXL} achieves significant improvements in average performance by using question-dependent prompts, when comparing the best variants using the (2)General and (2)Dependent templates, the latter obtains improvements of +3.5% on NExT-QA and +4.2% on STAR. Unsurprisingly, Q-ViD_{XXL} provides significant performance boosts when compared to its smaller version Q-ViD_{XL} achieving better performances on all type of questions in both datasets, showing a better capability to follow instructions, however, *this also demonstrates that using question-dependent frame captions to obtain specific information for the task at hand, performs better than general visual descriptions for zero-shot Video QA.*

Next, in Table 4 we investigate the impact on performance of the Question-Answering Instruction Prompt. We propose two variations that are shown

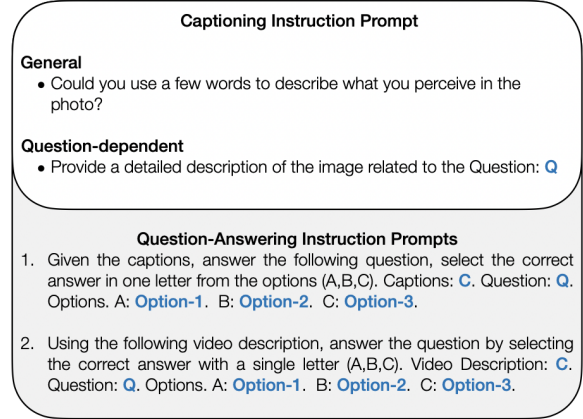


Figure 4: **Variation of QA prompt templates.** We focus on exploring two more complex and detailed variations for the Question Answering instruction prompt (Bottom). We use the best captioning templates (Top) for Q-ViD_{XL} (General) and Q-ViD_{XXL} (Dependent).

in Fig. 4 in addition to the Base QA prompt (Fig. 3). We use these new templates we aim to give more details to our reasoning module based on Flan-T5, because of this LLM is also a model trained with instructions, we explore if using more complex and detailed QA prompts we can achieve better performances. For this comparison we take the best variants (Table 3) of Q-ViD_{XL} and Q-ViD_{XXL} using the (2)General and (2)Dependent captioning prompts respectively for each model, and explore their performances with different QA instruction templates. As shown in Table 4 using more complex variants of the initial Base Question-Answering Instruction prompt does not have a big impact on performance of any of the models, it even slightly affects results in some cases, showing that the simplest base prompt was enough for the LLM to understand the task. With this ablation study we can highlight the fact that the input instructions used to obtain dedicated frame descriptions are far more important than elaborated question-answering instruction prompts for zero-shot video QA.

5 Conclusion

In this paper, we introduce Q-ViD, a simple, gradient-free approach for zero-shot video QA. Q-ViD turns video QA into textual QA using frame captions, for this it relies on an instruction-aware visual language model and uses question-dependent captioning instructions to obtain specific frame descriptions useful for the task at hand. This information is later used by a reasoning module with a question-answering instruction prompt to perform

Method	NExT-QA				STAR				
	Tem.	Cau.	Des.	Avg.	Int.	Seq.	Pre.	Fea.	Avg.
<i>Q-ViD_{XL}</i>									
(1) General	<u>57.3</u>	60.3	62.0	60.5	<u>47.0</u>	45.2	<u>42.7</u>	<u>42.2</u>	<u>44.3</u>
(2) General	57.8	<u>60.1</u>	60.8	<u>60.1</u>	47.4	<u>44.8</u>	44.7	42.8	44.9
(1) Dependent	55.9	59.8	57.5	59.1	45.0	41.7	40.5	40.2	41.8
(2) Dependent	56.6	58.8	<u>61.1</u>	59.0	45.8	40.6	40.2	39.5	41.5
<i>Q-ViD_{XXL}</i>									
(1) General	57.5	64.6	67.4	62.7	44.7	39.5	42.6	36.3	40.8
(2) General	57.1	64.8	68.0	62.8	44.6	39.5	<u>43.1</u>	38.7	41.5
(1) Dependent	62.0	<u>66.5</u>	<u>71.2</u>	<u>65.8</u>	<u>47.8</u>	<u>44.2</u>	42.1	<u>41.8</u>	<u>44.0</u>
(2) Dependent	<u>61.6</u>	67.6	72.2	66.3	48.2	47.2	43.9	43.4	45.7

Table 3: **Comparing the impact on performance using different captioning templates.** Base: Refer to the base captioning template (1) General: More detailed templates for general frame descriptions and (2) Dependent: templates for frame descriptions depending on a question. All experiments use the base QA instruction prompt.

Model	Templates		NExT-QA				STAR				
	Captioning	QA	Tem.	Cau.	Des.	Avg.	Int.	Seq.	Pre.	Fea.	Avg.
<i>Q-ViD_{XL}</i>	(2)General	Base	57.8	60.1	60.8	<u>60.1</u>	<u>47.4</u>	<u>44.8</u>	44.7	42.8	44.9
		(1)QA	56.4	<u>60.6</u>	<u>58.4</u>	60.2	47.7	44.9	<u>43.5</u>	<u>41.0</u>	<u>44.3</u>
		(2)QA	<u>56.8</u>	60.9	57.5	<u>60.1</u>	47.0	44.1	43.1	40.6	43.7
<i>Q-ViD_{XXL}</i>	(2)Dependent	Base	<u>61.6</u>	67.6	72.2	66.3	48.2	47.2	43.9	<u>43.4</u>	<u>45.7</u>
		(1)QA	61.7	<u>65.8</u>	<u>73.7</u>	<u>65.5</u>	<u>48.9</u>	<u>46.8</u>	<u>43.5</u>	43.8	45.8
		(2)QA	61.5	65.6	73.9	<u>65.5</u>	49.1	45.9	42.9	42.6	45.1

Table 4: **Performance using different variants for the QA template.** Base: Refer to the base QA instruction prompt. For the captioning prompts all models use their best variants, Q-ViD_{XL} with (2)General and Q-ViD_{XXL} with (2)Dependent. These results suggest that there is no improvements using more complex and detailed QA instruction prompts for the reasoning module, achieving more consistent performances with the simpler base template.

multiple-choice video QA. Our simple approach achieves competitive or even higher performances than more complex architectures and methods that rely on closed models like GPT’s. In our ablation studies we show that using dedicated instructions to get question-dependent captions work better than common prompts to get general descriptions from frames to perform video QA using captions.

Limitations

Even though, Q-ViD has shown to achieve strong performances for zero-shot video question answering, our approach suffers from some limitations. While the adopted instruction-aware multimodal model, InstructBlip, shows to successfully follow instructions from the question and extract meaningful information that can help the reasoning module to come up with the right answer, we have seen that in some cases the model tends to show hallucinations in the captions, or generate direct short one-word answers instead of a detailed and

question-specific description of the image. On the other hand, even though experiments with really long videos are not in the scope of this paper, our approach would no be recommended in those cases, due to the high memory usage that comes with saving detailed frame captions to create an entire video description, which would also affect the LLM-based reasoning module because of limited amount of tokens allowed as input or due to memory constrains to process the entire video description.

References

Jean-Baptiste Alayrac, Jeff Donahue, Pauline Luc, Antoine Miech, Iain Barr, Yana Hasson, Karel Lenc, Arthur Mensch, Katherine Millican, Malcolm Reynolds, Roman Ring, Eliza Rutherford, Serkan Cabi, Tengda Han, Zhitao Gong, Sina Samangooei, Marianne Monteiro, Jacob L Menick, Sebastian Borgeaud, Andy Brock, Aida Nematzadeh, Sahand Sharifzadeh, Mikołaj Bińkowski, Ricardo Barreira, Oriol Vinyals, Andrew Zisserman, and Karén Si-

573	monyan. 2022. Flamingo: a visual language model for few-shot learning . In <i>Advances in Neural Information Processing Systems</i> , volume 35, pages 23716–23736. Curran Associates, Inc.	630
574		631
575		632
576		
577	William Berrios, Gautam Mittal, Tristan Thrush, Douwe Kiela, and Amanpreet Singh. 2023. Towards language models that can see: Computer vision through the lens of natural language .	633
578		634
579		635
580		636
		637
		638
		639
581	Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel Ziegler, Jeffrey Wu, Clemens Winter, Chris Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei. 2020. Language models are few-shot learners . In <i>Advances in Neural Information Processing Systems</i> , volume 33, pages 1877–1901. Curran Associates, Inc.	640
582		641
583		642
584		643
585		
586		644
587		645
588		646
589		647
590		648
591		649
592		650
593		651
594		
595	Jun Chen, Deyao Zhu, Kilichbek Haydarov, Xiang Li, and Mohamed Elhoseiny. 2023a. Video chatcaptioner: Towards the enriched spatiotemporal descriptions. <i>arXiv preprint arXiv:2304.04227</i> .	652
596		653
597		654
598		655
		656
		657
599	Liangyu Chen, Bo Li, Sheng Shen, Jingkan Yang, Chunyuan Li, Kurt Keutzer, Trevor Darrell, and Zhiwei Liu. 2023b. Language models are visual reasoning coordinators . In <i>ICLR 2023 Workshop on Mathematical and Empirical Understanding of Foundation Models</i> .	658
600		659
601		660
602		661
603		662
604		663
605	Xi Chen, Xiao Wang, Soravit Changpinyo, AJ Piergiovanni, Piotr Padlewski, Daniel Salz, Sebastian Goodman, Adam Grycner, Basil Mustafa, Lucas Beyer, Alexander Kolesnikov, Joan Puigcerver, Nan Ding, Keran Rong, Hassan Akbari, Gaurav Mishra, Linting Xue, Ashish Thapliyal, James Bradbury, Weicheng Kuo, Mojtaba Seyedhosseini, Chao Jia, Burcu Karagol Ayan, Carlos Riquelme, Andreas Steiner, Anelia Angelova, Xiaohua Zhai, Neil Houlsby, and Radu Soricut. 2023c. Pali: A jointly-scaled multilingual language-image model .	664
606		665
607		666
608		667
609		668
610		669
611		670
612		671
613		672
614		673
615		
616	Hyung Won Chung, Le Hou, Shayne Longpre, Barret Zoph, Yi Tay, William Fedus, Yunxuan Li, Xuezhi Wang, Mostafa Dehghani, Siddhartha Brahma, Albert Webson, Shixiang Shane Gu, Zhuyun Dai, Mirac Suzgun, Xinyun Chen, Aakanksha Chowdhery, Alex Castro-Ros, Marie Pellat, Kevin Robinson, Dasha Valter, Sharan Narang, Gaurav Mishra, Adams Yu, Vincent Zhao, Yanping Huang, Andrew Dai, Hongkun Yu, Slav Petrov, Ed H. Chi, Jeff Dean, Jacob Devlin, Adam Roberts, Denny Zhou, Quoc V. Le, and Jason Wei. 2022. Scaling instruction-finetuned language models .	674
617		675
618		676
619		677
620		
621		678
622		679
623		680
624		681
625		
626		682
627		683
		684
628	Wenliang Dai, Junnan Li, Dongxu Li, Anthony Meng Huat Tiong, Junqi Zhao, Weisheng Wang,	685
629		
	Boyang Li, Pascale Fung, and Steven Hoi. 2023. Instructblip: Towards general-purpose vision-language models with instruction tuning .	
	Yuxin Fang, Wen Wang, Binhui Xie, Quan Sun, Ledell Wu, Xinggang Wang, Tiejun Huang, Xinlong Wang, and Yue Cao. 2023. Eva: Exploring the limits of masked visual representation learning at scale. In <i>Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)</i> , pages 19358–19369.	
	Taraneh Ghandi, Hamidreza Pourreza, and Hamidreza Mahyar. 2023. Deep learning approaches on image captioning: A review . <i>ACM Computing Surveys</i> , 56(3):1–39.	
	Deepanway Ghosal, Navonil Majumder, Roy Lee, Rada Mihalcea, and Soujanya Poria. 2023. Language guided visual question answering: Elevate your multimodal language model using knowledge-enriched prompts . In <i>Findings of the Association for Computational Linguistics: EMNLP 2023</i> , pages 12096–12102, Singapore. Association for Computational Linguistics.	
	Yash Goyal, Tejas Khot, Douglas Summers-Stay, Dhruv Batra, and Devi Parikh. 2017. Making the v in vqa matter: Elevating the role of image understanding in visual question answering. In <i>Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)</i> .	
	Jiaxian Guo, Junnan Li, Dongxu Li, Anthony Meng Huat Tiong, Boyang Li, Dacheng Tao, and Steven Hoi. 2023. From images to textual prompts: Zero-shot visual question answering with frozen large language models. In <i>Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)</i> , pages 10867–10877.	
	Jordan Hoffmann, Sebastian Borgeaud, Arthur Mensch, Elena Buchatskaya, Trevor Cai, Eliza Rutherford, Diego de Las Casas, Lisa Anne Hendricks, Johannes Welbl, Aidan Clark, Tom Hennigan, Eric Noland, Katie Millican, George van den Driessche, Bogdan Damoc, Aurelia Guy, Simon Osindero, Karen Simonyan, Erich Elsen, Jack W. Rae, Oriol Vinyals, and Laurent Sifre. 2022. Training compute-optimal large language models .	
	Neil Houlsby, Andrei Giurgiu, Stanislaw Jastrzebski, Bruna Morrone, Quentin de Laroussilhe, Andrea Gesmundo, Mona Attariyan, and Sylvain Gelly. 2019. Parameter-efficient transfer learning for nlp .	
	Yushi* Hu, Hang* Hua, Zhengyuan Yang, Weijia Shi, Noah A Smith, and Jiebo Luo. 2022. Promptcap: Prompt-guided task-aware image captioning. <i>arXiv preprint arXiv:2211.09699</i> .	
	Dohwan Ko, Ji Soo Lee, Wooyoung Kang, Byungseok Roh, and Hyunwoo J Kim. 2023. Large language models are temporal and causal reasoners for video question answering. In <i>EMNLP</i> .	

686	Jie Lei, Licheng Yu, Mohit Bansal, and Tamara L Berg.	Sasank Chilamkurthy, Benoit Steiner, Lu Fang, Jun-	740
687	2018. Tvqa: Localized, compositional video ques-	jie Bai, and Soumith Chintala. 2019. Pytorch: An	741
688	tion answering. In <i>EMNLP</i> .	imperative style, high-performance deep learning li-	742
		brary .	743
689	Dongxu Li, Junnan Li, Hung Le, Guangsen Wang, Sil-	Archiki Prasad, Elias Stengel-Eskin, and Mohit Bansal.	744
690	vio Savarese, and Steven C. H. Hoi. 2022. Lavis: A	2023. Rephrase, augment, reason: Visual grounding	745
691	library for language-vision intelligence .	of questions for vision-language models.	746
692	Jiapeng Li, Ping Wei, Wenjuan Han, and Lifeng Fan.	Dustin Schwenk, Apoorv Khandelwal, Christopher	747
693	2023a. Intentqa: Context-aware video intent rea-	Clark, Kenneth Marino, and Roozbeh Mottaghi. 2022.	748
694	soning. In <i>Proceedings of the IEEE/CVF Interna-</i>	A-okvqa: A benchmark for visual question answer-	749
695	<i>tional Conference on Computer Vision (ICCV)</i> , pages	ing using world knowledge .	750
696	11963–11974.		
697	Junnan Li, Dongxu Li, Silvio Savarese, and Steven Hoi.	Dídac Surís, Sachit Menon, and Carl Vondrick. 2023.	751
698	2023b. Blip-2: Bootstrapping language-image pre-	Vipergpt: Visual inference via python execution for	752
699	training with frozen image encoders and large lan-	reasoning .	753
700	guage models .		
701	Kunchang Li, Yali Wang, Yinan He, Yizhuo Li,	Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier	754
702	Yi Wang, Yi Liu, Zun Wang, Jilan Xu, Guo	Martinet, Marie-Anne Lachaux, Timothée Lacroix,	755
703	Chen, Ping Luo, Limin Wang, and Yu Qiao. 2024.	Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal	756
704	Mvbench: A comprehensive multi-modal video un-	Azhar, Aurelien Rodriguez, Armand Joulin, Edouard	757
705	derstanding benchmark .	Grave, and Guillaume Lample. 2023. Llama: Open	758
		and efficient foundation language models .	759
706	Linjie Li, Yen-Chun Chen, Yu Cheng, Zhe Gan, Licheng	Oriol Vinyals, Alexander Toshev, Samy Bengio, and	760
707	Yu, and Jingjing Liu. 2020a. Hero: Hierarchical	Dumitru Erhan. 2015. Show and tell: A neural image	761
708	encoder for video+ language omni-representation pre-	caption generator. In <i>Proceedings of the IEEE Con-</i>	762
709	training. In <i>EMNLP</i> .	<i>ference on Computer Vision and Pattern Recognition</i>	763
		(<i>CVPR</i>).	764
710	Linjie Li, Yen-Chun Chen, Yu Cheng, Zhe Gan, Licheng	Junke Wang, Dongdong Chen, Chong Luo, Xiyang Dai,	765
711	Yu, and Jingjing Liu. 2020b. HERO: Hierarchical	Lu Yuan, Zuxuan Wu, and Yu-Gang Jiang. 2023.	766
712	encoder for Video+Language omni-representation	Chatvideo: A tracklet-centric multimodal and versa-	767
713	pre-training . In <i>Proceedings of the 2020 Conference</i>	tile video understanding system .	768
714	<i>on Empirical Methods in Natural Language Process-</i>		
715	<i>ing (EMNLP)</i> , pages 2046–2065, Online. Association	Yi Wang, Kunchang Li, Yizhuo Li, Yinan He, Bingkun	769
716	for Computational Linguistics.	Huang, Zhiyu Zhao, Hongjie Zhang, Jilan Xu, Yi Liu,	770
		Zun Wang, Sen Xing, Guo Chen, Juntong Pan, Ji-	771
717	Yicong Li, Junbin Xiao, Chun Feng, Xiang Wang, and	ashuo Yu, Yali Wang, Limin Wang, and Yu Qiao.	772
718	Tat-Seng Chua. 2023c. Discovering spatio-temporal	2022a. Internvideo: General video foundation mod-	773
719	rationales for video question answering .	els via generative and discriminative learning .	774
720	Haotian Liu, Chunyuan Li, Yuheng Li, and Yong Jae	Zhenhailong Wang, Manling Li, Ruochen Xu, Lu-	775
721	Lee. 2023. Improved baselines with visual instruc-	wei Zhou, Jie Lei, Xudong Lin, Shuhang Wang,	776
722	tion tuning .	Ziyi Yang, Chenguang Zhu, Derek Hoiem, Shih-Fu	777
		Chang, Mohit Bansal, and Heng Ji. 2022b. Language	778
723	Kenneth Marino, Mohammad Rastegari, Ali Farhadi,	models with image descriptors are strong few-shot	779
724	and Roozbeh Mottaghi. 2019. Ok-vqa: A visual	video-language learners .	780
725	question answering benchmark requiring external		
726	knowledge. In <i>Proceedings of the IEEE/CVF Con-</i>	Bo Wu, Shoubin Yu, Zhenfang Chen, Joshua B. Tenen-	781
727	<i>ference on Computer Vision and Pattern Recognition</i>	baum, and Chuang Gan. 2021. STAR: A benchmark	782
728	(<i>CVPR</i>).	for situated reasoning in real-world videos . In <i>Thirty-</i>	783
		<i>fifth Conference on Neural Information Processing</i>	784
729	Liliane Momeni, Mathilde Caron, Arsha Nagrani,	<i>Systems Datasets and Benchmarks Track (Round 2)</i> .	785
730	Andrew Zisserman, and Cordelia Schmid. 2023.	Junbin Xiao, Xindi Shang, Angela Yao, and Tat-Seng	786
731	Verbs in action: Improving verb understanding in	Chua. 2021. Next-qa: Next phase of question-	787
732	video-language models. In <i>Proceedings of the</i>	answering to explaining temporal actions. In <i>Pro-</i>	788
733	<i>IEEE/CVF International Conference on Computer</i>	<i>ceedings of the IEEE/CVF Conference on Computer</i>	789
734	<i>Vision (ICCV)</i> , pages 15579–15591.	<i>Vision and Pattern Recognition (CVPR)</i> , pages 9777–	790
		9786.	791
735	Adam Paszke, Sam Gross, Francisco Massa, Adam	Antoine Yang, Antoine Miech, Josef Sivic, Ivan Laptev,	792
736	Lerer, James Bradbury, Gregory Chanan, Trevor	and Cordelia Schmid. 2022. Zero-shot video ques-	793
737	Killeen, Zeming Lin, Natalia Gimelshein, Luca	tion answering via frozen bidirectional language mod-	794
738	Antiga, Alban Desmaison, Andreas Köpf, Edward	els. In <i>NeurIPS</i> .	795
739	Yang, Zach DeVito, Martin Raison, Alykhan Tejani,		

Qinghao Ye, Haiyang Xu, Guohai Xu, Jiabo Ye, Ming Yan, Yiyang Zhou, Junyang Wang, Anwen Hu, Pengcheng Shi, Yaya Shi, Chaoya Jiang, Chenliang Li, Yuanhong Xu, Hehong Chen, Junfeng Tian, Qian Qi, Ji Zhang, and Fei Huang. 2023. [mplug-owl: Modularization empowers large language models with multimodality](#).

Shoubin Yu, Jaemin Cho, Prateek Yadav, and Mohit Bansal. 2023. Self-chained image-language model for video localization and question answering. In *NeurIPS*.

Andy Zeng, Maria Attarian, Brian Ichter, Krzysztof Choromanski, Adrian Wong, Stefan Welker, Federico Tombari, Aveek Purohit, Michael Ryoo, Vikas Sindhwani, Johnny Lee, Vincent Vanhoucke, and Pete Florence. 2022. [Socratic models: Composing zero-shot multimodal reasoning with language](#).

Ce Zhang, Taixi Lu, Md Mohaiminul Islam, Ziyang Wang, Shoubin Yu, Mohit Bansal, and Gedas Bertasius. 2023a. [A simple llm framework for long-range video question-answering](#).

Renrui Zhang, Jiaming Han, Chris Liu, Peng Gao, Aojun Zhou, Xiangfei Hu, Shilin Yan, Pan Lu, Hongsheng Li, and Yu Qiao. 2023b. [Llama-adapter: Efficient fine-tuning of language models with zero-init attention](#).

Yaoyao Zhong, Wei Ji, Junbin Xiao, Yicong Li, Weihong Deng, and Tat-Seng Chua. 2022. [Video question answering: Datasets, algorithms and challenges](#). In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pages 6439–6455, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.

Deyao Zhu, Jun Chen, Xiaoqian Shen, Xiang Li, and Mohamed Elhoseiny. 2023. Minigt-4: Enhancing vision-language understanding with advanced large language models. *arXiv preprint arXiv:2304.10592*.

A Licences

We use standard licenses from the community for the datasets, codes, and models that we used in this paper:

- **NExT-QA** (Xiao et al., 2021): MIT
- **STAR** (Wu et al., 2021): Apache
- **How2QA** (Li et al., 2020a): MIT
- **TVQA** (Lei et al., 2018): MIT
- **IntentQA** (Li et al., 2023a): N/A
- **SeViLa** (Yu et al., 2023): BSD 3 - Clause
- **LAVIS** (Li et al., 2022): BSD 3-Clause

• **Pytorch** (Paszke et al., 2019): BSD Style

• **Q-ViD** (Ours): Available soon under the BSD Style license.

B Use of Artifacts

In this work we adopt a open multimodal model, InstructBLIP (Dai et al., 2023), its application in our approach is consistent with its original intended use. For Q-ViD we will release our code and we hope will be useful for future works.