

VELOCITI: Benchmarking Video-Language Compositional Reasoning with Strict Entailment

Darshana Saravanan¹Varun Gupta¹Darshan Singh¹Zeeshan Khan²Vineet Gandhi¹Makarand Tapaswi¹¹CVIT, IIT Hyderabad, India²Inria, Paris, France [katha-ai.github.io/projects/velociti](https://github.com/katha-ai/projects/velociti)

Abstract

A fundamental aspect of compositional reasoning in a video is associating people and their actions across time. Recent years have seen great progress in general-purpose vision/video models and a move towards long-video understanding. While exciting, we take a step back and ask: are today’s models good at compositional reasoning on short videos? To this end, we introduce VELOCITI, a benchmark to study Video-LLMs by disentangling and assessing the comprehension of agents, actions, and their associations across multiple events. We adopt the Video-Language Entailment setup and propose StrictVLE that requires correct classification (rather than ranking) of the positive and negative caption. We evaluate several models and observe that even the best, LLaVA-OneVision (44.5%) and Gemini-1.5-Pro (49.3%), are far from human accuracy at 93.0%. Results show that action understanding lags behind agents, and negative captions created using entities appearing in the video perform worse than those obtained from pure text manipulation. We also present challenges with ClassicVLE and multiple-choice (MC) evaluation, strengthening our preference for StrictVLE. Finally, we validate that our benchmark requires visual inputs of multiple frames making it ideal to study video-language compositional reasoning.

1. Introduction

Near a parking lot, a man in a black hat smiles in a friendly way at a woman in a purple shirt. To a reader, this dense description paints a clear picture about a short snippet (event) of a video clip. We build a mental model of two people (referred here by their clothing), at a specified location, and a short interaction between them. Reading further, *the woman claps as a man in grey pants spins on one leg*. We are able to associate that it is the same woman who is now cheering at a third person (likely) that is performing stunts.

The above example illustrates an intelligent agent’s abil-

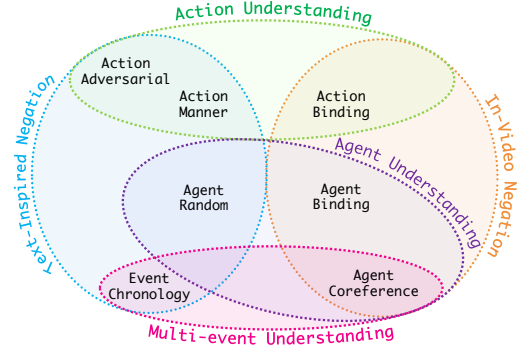


Figure 1. A Venn diagram grouping VELOCITI’s seven tests (in black) that evaluate a Video-LLM across different facets: **Agent Understanding**, **Action Understanding**, and **Multi-event Understanding**. The benchmark is formulated as video-language entailment, where negative captions are created by manipulating text (**Text-inspired Negation**) or from other parts of the same video (**In-Video Negation**). Best seen in color.

ity to perform compositional reasoning. For video-language models, we scope this in two steps: (i) comprehend atomic entities, *e.g.* *people* and *actions*; and (ii) reason about them compositionally and across time by building associations¹.

In recent years, strong visual (image) encoders are combined with powerful Large Language Models (LLMs) to advance general-purpose vision [5, 9, 11, 23, 45]. A similar approach is adapted for videos to create Video-LLMs [23, 28, 45, 48]. Keeping pace with the development of new models, there is a flurry of work on evaluating them (Tab. 1). Video researchers are also creating benchmarks to study long video comprehension [7, 13, 15, 37]. However, we take a step back and ask, *are today’s Video-LLMs ready to take on such challenges?* Specifically, are they good at compositional reasoning in short videos, arguably a prerequisite to tackle complex and long videos?

To this end, we introduce the *VELOCITI*, a benchmark

¹ Associations can be thought as *implicit* tuples that a model attempts to build while watching a video. Some examples include *person-attribute* tuples: (man1, black hat), (woman, purple shirt), (man2, grey pants); *agent-action* tuples: (man1, smiles at, woman), (man2, spins on one leg), (woman, claps at, man2); or *action-manner* tuples: (smile, friendly way).

 Equal contribution

that studies Video et Language Compositionality through Time. We adopt the video-language entailment (VLE) evaluation setup [4] where a model is prompted to predict whether a video entails a caption (‘Yes’ for an aligned or positive caption and ‘No’ for a misaligned or negative caption). Through a suite of seven tests, we are able to disentangle and assess a model’s ability to comprehend *agents*, *actions*, and their associations across *multiple events* through time. As illustrated in Fig. 1, we group the 7 tests based on: (i) the specific facet of a model’s ability (agent, action, multi-event), or (ii) the strategy used to create the negative caption (*Text-Inspired Negation* vs. *In-Video Negation*). Note that although the tests (Sec. 3.2) have varying levels of difficulty, they are all important as they shed light on whether a model is able to solve a specific facet of video-language compositional reasoning.

Our videos are sourced from the VidSitu dataset and are accompanied by action and semantic role label (SRL) annotations for multiple events in a short movie clip [39]. The videos are diverse and feature multiple agents and actions across complex editing and shot changes, while dense SRL succinctly describes *who* did *what* with/to *whom*, *where*, and (sometimes) *how*. Importantly, each SRL only describes a single event, requiring models to implicitly localize the event in the video before solving the test.

Strict entailment. In the classic VLE setup, benchmarks typically check if the entailment score for the positive caption is higher than the negative caption [25, 40, 47]. While this traces back to visual-semantic embedding models [12, 14, 35], it is unsuitable for evaluating modern Video-LLMs that generate text (and not similarity scores).

We propose a strict entailment scoring mechanism where Video-LLMs should output ‘Yes’ for an aligned caption and ‘No’ for a misaligned one. Our analysis reveals that models produce marginally different entailment scores for the positive and negative captions attaining good performance on ClassicVLE, but predict ‘Yes’ for both. This is critical as VLE evaluates a model’s ability to reject (partially) misaligned descriptions. Poor performance here implies that the model may produce erroneous outputs on other tasks (e.g. question-answering) and assumes that partial hallucinations (like negative captions) are acceptable.

Contributions. We summarize our contributions and findings below: (i) We propose VELOCITI, a new benchmark that evaluates compositional reasoning of video-language models. Our test suite sheds light on a model’s ability to perceive and reason about *agents* and *actions* across *multiple events*, identifying challenges for improvement (Sec. 3). (ii) We propose a strict metric for video-language entailment that requires a model to produce ‘Yes’ for an aligned caption and ‘No’ for the corresponding misaligned caption (Sec. 4). (iii) We evaluate both open and closed models and show that they struggle with compositional rea-

soning. While larger models such as LLaVA-OneVision-72B (OV-72B) [23] tend to perform better than smaller ones (OV-7B), even the best commercial model (Gemini-1.5-Pro [10]) achieves 49.3% accuracy, about half that of humans at 93.0%. (iv) Our experiments reveal important findings: a) Understanding actions is harder than agents for open models, and b) tests incorporating in-video negation are more challenging than text-inspired negation (Sec. 5.1). c) Smaller models are predisposed towards ‘Yes’ for the entailment task (Sec. 5.2). d) ClassicVLE hides information as entailment scores of positive and negative captions are often close to each other, likely due to subtle differences between them (Sec. 5.3). e) Multiple-choice (MC) evaluation is unsuitable due to a choice bias observed even in large and closed models (Sec. 5.4). f) Finally, we show that VELOCITI requires visual inputs and multiple frames and cannot be solved with text-only or single-frame models (Sec. 5.5).

2. Related Work

Several benchmarks exist to evaluate image-language compositionality (Winoground [42], COLA [38], MMVP [44], and others [17, 26, 32, 49, 52]). They require identifying the correct caption among distractors, exposing models failure to bind concepts [20]. We focus on short complex videos.

Video-language benchmarks broadly related to our work are presented in Tab. 1. We discuss differences to closely related work here. Among previous benchmarks that study compositional reasoning, our work differs due to the (i) emphasis on a test suite that provides disentangled understanding of agents and actions across multiple events; (ii) task formulation as *strict* video-language entailment (unlike TestOfTime [3], VideoCon [4], VITATECS [25], Vinoground [51]); (iii) explicit use of text-inspired *and* in-video negation (unlike TestOfTime [3], VideoCon [4], MV-Bench [24], VITATECS [25]); and (iv) use of short complex videos (e.g. compared to indoor, single-agent Charades [41] in AGQA [16], STAR [46]).

While contrast captions are a popular strategy [3, 4, 8, 24, 25, 33, 51], the structured SRL annotations used in VELOCITI facilitate evaluating specific aspects of a model’s capabilities. Further, different from comprehensive benchmarks that evaluate holistic video understanding (e.g. SEED-Bench-2 [21], CVRR-ES [19], Video-MME [15]), we focus on the fundamental ability of compositional understanding and highlight major shortcomings. Importantly, our tests are designed to prevent text-only and single-frame models from solving them (validated empirically), guarding against issues highlighted by ATP [6] and recently TVBench [8].

Video-Language Entailment (VLE) is posed as a binary classification task [40]. Given a premise (the video) and a hypothesis (the caption), a model should determine if the

Benchmark	Task Setup	Comp	In-V Neg	Strict VLE	Test Creation	Human Eval	Video Duration	Domain (Source)
AGQA [16] CVPR’21	OQA, MCQ	✓	✗	NA	T, SG	✓	30s	Open (ActionGenome, Charades)
STAR [46] NeurIPSDB’21	MCQ	✓	✓	NA	H, T, SG	✗	30s	Indoor (Charades)
ContrastSets [33] NAACL’22	MCQ	✓	✗	NA	H, T, LLM	✓	-	Mixed (MSR-VTT, LSMDC)
TestOfTime [3] CVPR’23	E	✗	✗	✗	T	✗	5-30s	Open (TEMPO, ANet Cap., Charades)
Perception Test [34] NeurIPSDB’23	MCQ	✗	✗	NA	H	✓	23s	Indoor (Manual)
Cinepile [37] CVPRW’24	MCQ	✗	✗	NA	H, GPT-4	✓	2-3m	Movies (MovieClips channel)
VideoCon [4] CVPR’24	E, OQA	✓	✗	✗	H, PaLM-2	✗	10-30s	Open (MSR-VTT, VATEX, TEMPO)
SEED-Bench-2 [21] CVPR’24	MCQ	✗	✗	NA	H, GPT-4	✗	-	Open (Charades, SSV2, EK100)
MV-Bench [24] CVPR’24	MCQ	✓	✗	NA	T, ChatGPT	✗	5-35s	Mixed (Charades-STA, MoVQA, +9)
TempCompass [31] ACLFindings’24	MCQ, E, VC	✓	✗	✗	H, GPT-3.5	✓	30s	Open (Shutterstock)
MMBench-Video [13] NeurIPSDB’24	OQA	✗	✗	NA	H	✗	30s-6m	Open (YouTube)
VITATECS [25] ECCV’24	E	✓	✗	✗	H, GPT-3.5	✓	10s	Open (MSRVTT, VATEX)
CVRR-ES [19] arXiv-2405	OQA	✗	✗	NA	H, GPT-3.5	✓	2-183s	Open (SSV2, CATER, +5)
Video-MME [15] arXiv-2405	MCQ	✗	✓	NA	H	✗	11s-1h	Open (YouTube)
VideoVista [27] arXiv-2406	MCQ	✗	✗	NA	T, GPT-4, GPT-4o	✗	131s	Mixed (Panda-70M)
Vinoground [51] arXiv-2410	E	✓	✗	✗	H, GPT-4	✓	10s	Open (VATEX)
TVBench [8] arXiv-2410	MCQ	✓	✗	NA	T	✗	-	Mixed (STAR, CLEVRER, +6)
VELOCITI (Ours)	E	✓	✓	✓	H, T, LLM	✓	10s	Movies (VidSitu)

Table 1. We review video-language benchmarks and highlight key differences to VELOCITI. Benchmarks use various ‘Task Setups’: Entailment (E), Multiple Choice (MCQ), Open-ended Question-Answering (OQA), and Video Captioning (VC). We compare VELOCITI against benchmarks that test Compositionality (‘Comp’) or have In-Video Negation (‘In-V Neg’). In ‘StrictVLE’, benchmarks not adopting VLE are marked not applicable (NA). Acronyms in the ‘Test Creation’ column are: template (T), scene graph (SG), open large language model (LLM), and human (H). The ‘Domains’ are of 3 types: Open (natural videos), Movies, and Mixed (natural & movies). Different from others, VELOCITI introduces StrictVLE and features tests with negative captions created from entities appearing in the same video.

hypothesis logically follows (entails) from the premise. Entailment was first used with images in [47] and adopted by [3, 4, 25] for videos. Given the rise of Vision LLMs, entailment scores are computed using the likelihood over specific words in the vocabulary [4, 22, 29]. However, most works only require that the positive caption scores higher than the negative [25, 47], or with a margin [22]. We propose a more demanding form, *StrictVLE*, that unlike ClassicVLE, is applicable to both open and closed models (without likelihood scores). Specifically, we independently require that the model entails the positive caption *and* does not entail the negative caption. While this looks simple, we find that models do not sufficiently distinguish positive and negative captions and tend to answer ‘Yes’ for both.

3. VELOCITI Benchmark

We evaluate compositional reasoning using dynamic 10 s movie clips and SRL annotations from the VidSitu dataset. We propose *seven tests* to evaluate model’s comprehension of agents and actions across multiple events through time. Each test consists of $\{V, C^+, C^-\}$: video clip V , a positive caption C^+ that is aligned with a part of the video, and a negative caption C^- that is *not* aligned to the video. We require models to *independently* assess each caption and classify them as V entails C^+ *and* V does not entail C^- .

3.1. From SRL to a Video-Caption Pair

In VidSitu, videos are divided into five 2 s events [39] (total 10 s duration). Each event is annotated by the most salient

action and the corresponding SRL capturing: *who* is doing the action (agent), *with / to whom* (patient or receiver), *with what* (instrument, if applicable), *where* (scene or location), *how* (manner or adverb), and *why* (purpose).

We use an open LLM (LLaMA-3 [2]) to convert the structured SRL dictionary of each event into a caption. The LLM is prompted to combine the atomic concepts into a fluent caption (prompt in Fig. 10). We filter 864 videos from the validation set and generate 3101 high-quality (V, C) pairs with captions that are faithful to the SRL annotations. Depending on the test, these captions are directly used as C^+ or used to *form* C^+ .

Note, movie events are not bounded by 2 s intervals and the SRL annotations may spill into the neighboring events. Thus, we make a conscious choice to pair the caption C with the entire 10 s video V . To correctly decide whether V entails C , a model needs to implicitly *localize* to the appropriate temporal region in the video. This prevents a single frame bias as reported by Atemporal Probe in [6].

3.2. VELOCITI Tests

We motivate and describe the seven tests below. Fig. 2 shows an example of each test grouped based on the process used to create C^- : (i) *text-inspired negation* typically creates C^- without looking at the video; and (ii) *in-video negation*, a key contribution of our work, uses a different entity appearing in the same video to create C^- . Both are important as they help us identify pitfalls of current models.

0. Control Test. We start with a control test to establish a

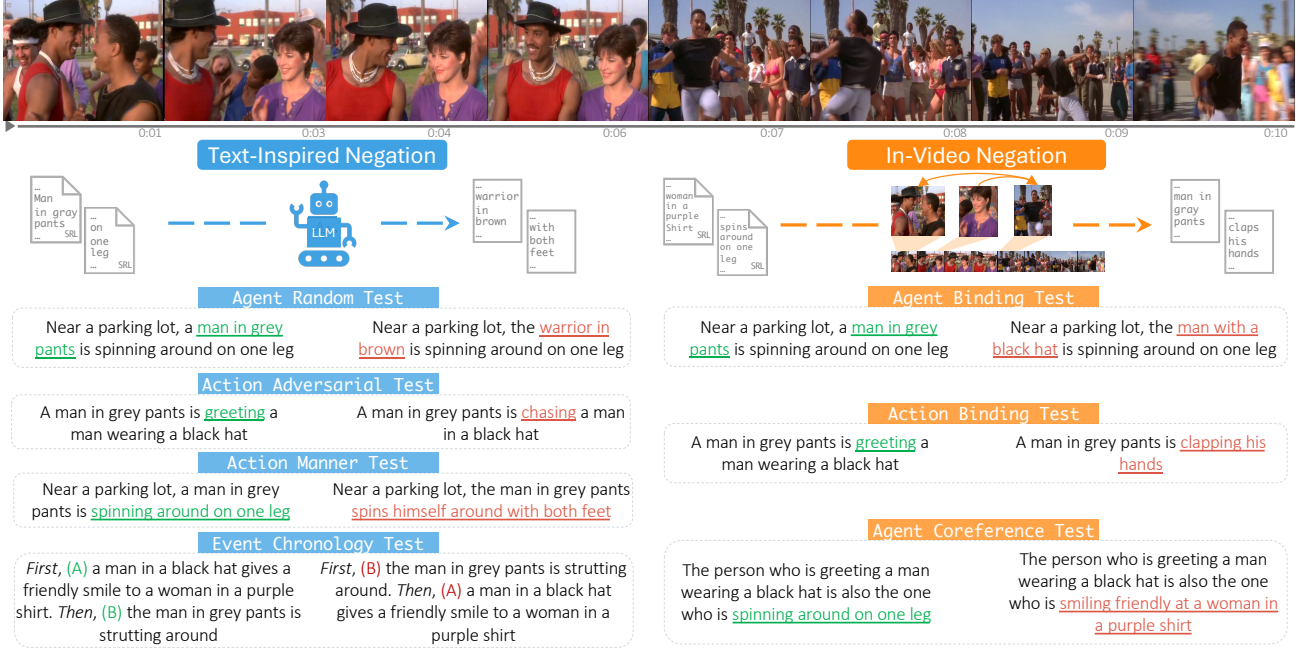


Figure 2. VELOCITY evaluates Video-LLMs’ video-language entailment capabilities on complex movie clips with dense semantic role label (SRL) annotations from the VidSitu dataset [39]. Positive and negative captions are shown side-by-side for each test with the key difference highlighted with green/red. Negative captions are created by (i) manipulating text using an LLM (**Text-Inspired Negation**) or (ii) replacing agents or actions by others that appear in the same video (**In-Video Negation**). We also demonstrate how the same positive caption can be used to create negative captions differently (see Agent Random vs. Agent Binding test; or Action Adversarial vs. Action Binding test). Each test evaluates models for different facets of compositional reasoning as described in Sec. 3.2. The 10 s video clip used in this example can be viewed here: <https://www.youtube.com/embed/bt6-F11LZsQ?start=25&end=35>.

baseline understanding. Here, C^+ is as described in Sec. 3.1 and C^- is simply a positive caption of some other random video, making it easily discernible.

1. Agent Random Test. C^- is created by replacing the correct *agent* with another agent that does not appear in the video (C^+ is as above). Solving this test requires a model to implicitly localize the event based on the action and identify who is present/absent in the video. We ensure that the replacing agent is not a hypernym (e.g., “man in a shirt” is not replaced by “man”). The SRL dictionary is updated with the random agent and the LLM generates C^- .

2. Agent Binding Test also replaces the agent. Different from above, the replaced agent is chosen from the *same video* making it an **in-video negation**. This subtle difference requires models to identify the correct agent and bind or associate it with the event description. Models cannot rely purely on presence/absence to solve this task. Fig. 2 shows how the same C^+ can be modified to create both agent tests. Similar to above, the LLM generates C^- .

3. Agent Coreference Test. Coreference groups two or more phrases that refer to the same entity [18]. In a video, an agent can be referred to by their actions, e.g. in Fig. 2, the agent: *man in grey pants* is referred by: *the person who is (i) greeting a man wearing a black hat* or (ii) *spinning*

around on one leg. To create this test, we identify videos with the same person acting in two or more events and construct two references for that person. C^+ is formed by combining the referring expressions of the *same* agent, while C^- combines referring expressions of *different* agents. This test also features complex **in-video negation** as all concepts mentioned in both C^+ and C^- appear in the video. The captions are created using the template: *The person who is [Event A] is also the one who is [Event B]*. Solving this test requires models to associate the correct interactions of an agent across two events. Since the agent description is masked by *the person*, a model requires multi-level compositional reasoning, making this test particularly challenging.

4. Action Adversarial Test. C^+ is as described in Sec. 3.1 and C^- is created by replacing the *action* with an adversarial alternative (a plausible action determined through the text description) that does not appear in the video. Solving this test requires identifying the action that the agent is performing. Given the SRL dictionary, the LLM is prompted to first generate the adversarial action followed by C^- .

5. Action Manner Test typically features a C^+ that includes an adverb, emotion, or facial expression. C^- is generated by replacing this manner with a contradictory yet plausible alternative. Solving this test is challenging as it requires understanding subtle variations in an action. Simi-

lar to the test above, the LLM is prompted to first generate the contrasting manner followed by C^- .

6. Action Binding Test. Here, C^- is created by retaining the agent from C^+ and swapping the action and its modifiers with those from a different event within the *same video*. Solving this test requires models to localize events where the agent appears and bind them with the correct action. This is another test with **in-video negation** as actions described in both C^+ and C^- appear in the video. To create C^- , we identify an event in the same video with a different action performed by a different agent. Next, we replace the SRL dictionary of C^+ with the action (and relevant modifiers) and prompt the LLM to generate C^- .

7. Event Chronology Test. Our final test studies a model’s ability to confirm whether the video and caption follow the same event progression. Multiple events descriptions can be related through time using *before*, *after*, *first*, *then* [3]. C^+ is created by concatenating event captions (from Sec. 3.1) with the template “*First, [Event A]. Then, [Event B].*” where event A *precedes* B. C^- simply reverses them to “*First, [Event B]. Then, [Event A].*” The events are sampled at least 2s apart to prevent a chance of overlap.

Quality control. All test samples in VELOCITI are verified by humans to ensure that C^+ aligns with the video and C^- is misaligned. The number of samples in each test and other details are presented in Tab. 11.

4. StrictVLE Evaluation Metric

We adopt Video-Language Entailment (VLE) as the evaluation scheme for VELOCITI. Given an instruction I containing a video V and a caption C , model M is prompted to answer whether the video entails the caption through ‘Yes’/‘No’. We define the entailment score similar to [40]:

$$e(V, C) = \frac{p_M(\text{‘Yes’} | I(V, C))}{p_M(\text{‘Yes’} | I(V, C)) + p_M(\text{‘No’} | I(V, C))}, \quad (1)$$

where p_M denotes the model’s probability distribution over the entire vocabulary.

ClassicVLE [25, 47] considers that a model is correct when $e(V, C^+) > e(V, C^-)$. The random accuracy is 50%.

Narrative example. Consider a simple video of a red traffic light. Let C^+ , “The traffic light is red” score $p(\text{‘Yes’})=0.7$, $p(\text{‘No’})=0.3$; and C^- , “The traffic light is green” score $p(\text{‘Yes’})=0.6$, $p(\text{‘No’})=0.4$. As $e(V, C^+) > e(V, C^-)$, ClassicVLE considers this as a correct prediction. However, with greedy decoding (constrained to ‘Yes’ and ‘No’), the model will predict ‘Yes’ for C^- , which is objectively incorrect, and in this example scenario, dangerous.

StrictVLE. As models improve, it is important for our community to hold them to higher standards. We argue that relative ordering of the entailment scores is insufficient and we propose StrictVLE that requires models

to predict ‘Yes’ for C^+ and ‘No’ for the corresponding C^- . Specifically, StrictVLE considers a sample correct iff $e(V, C^+) > 0.5 \wedge e(V, C^-) < 0.5$ ² and has a random chance accuracy of 25%. The threshold 0.5 arises naturally, and is equivalent to greedy decoding of ‘Yes’/‘No’ in the response. This equivalence means StrictVLE also works on closed models without access to p_M .

Relation to multiple-choice (MC). While VLE performs *independent* evaluation of C^+ and C^- , MC provides *both* captions to the model at once. Seeing both captions makes the task easier as the model needs to predict the *more likely* option rather than independently assess each one (similar to ClassicVLE). In Sec. 5.4 we reveal biases of MC evaluation and show that our proposed StrictVLE is preferable.

5. Results and Discussion

We evaluate open and closed Video-LLMs on VELOCITI. First, we present results with StrictVLE, our primary evaluation strategy, and analyze entailment scores. We also compare results with ClassicVLE and MC, discussing some evaluation pitfalls. Finally, we evaluate blind and single-frame models to check for bias in the benchmark.

Models. We evaluate multiple open Video-LLMs: PLLaVA [48], Video-LLaVA (V-LLaVA) [28], Owl-Con [4], Qwen2-VL (QVL) [45], and LLaVA-OneVision (OV) [23]; closed models Gemini-1.5-Flash (Gem-1.5F), Gemini-1.5-Pro (Gem-1.5P) [10], GPT-4o [1]; and humans. Due to compute and cost constraints, we evaluate QVL-72B (at native video resolution), closed models, and humans on a subset of 150 samples from each test (created once by random selection). More details in Appendix A.6.

5.1. Evaluation with StrictVLE

We report model performance in Tab. 2 and discuss various facets of model understanding highlighted in Fig. 1 (agents, actions, multiple events, and negation strategies).

Control vs. VELOCITI average. Old open models (P-LLaVA, Owl-Con) struggle on the StrictVLE setup as they have a strong bias to predict ‘Yes’. This leads to poor performance on the control test and the benchmark average (first and last column). V-LLaVA and subsequent models, QVL and OV, obtain decent accuracies on the control test, ranging from 65-85%. Compared to the control tests, performance dips strongly on the benchmark average, with the best model, OV-72B obtaining 43.2% accuracy (36.2% lower than control). On the benchmark subset, while GPT-4o posts a 46.2% accuracy on average, it performs poorly on the control tests (analysis in Sec. 5.2). Inversely, Gem-1.5F achieves 91.9% on the control tests, but is close to random on the benchmark (23.9%). The best model, Gem-1.5P,

²Note, this is different from the Winoground [42] setup that has: (i) 2 images/videos and 2 captions; and (ii) still uses relative scoring.

Model	Ctrl	Ag Rand	Ag Bind	Ag Cref	Act Adv	Act Man	Act Bind	Ev Chr	Avg
Random	25.0	25.0	25.0	25.0	25.0	25.0	25.0	25.0	25.0
P-LLaVA	1.3	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
Owl-Con	24.3	3.4	0.7	0.0	4.3	2.8	0.6	0.1	1.7
V-LLaVA	65.8	16.4	7.6	0.3	8.7	3.3	10.6	3.9	7.3
QVL-7B	84.6	39.1	13.5	6.5	17.8	17.5	16.4	0.4	15.9
OV-7B	81.6	56.7	32.9	8.0	29.7	30.6	36.4	30.5	32.1
OV-72B	79.3	63.7	45.4	38.6	33.1	29.3	45.1	46.5	43.1
VELOCITI Subset									
Gem-1.5F	91.9	56.4	23.8	4.7	32.9	21.6	25.0	2.7	23.9
QVL-72B	82.7	56.0	29.3	35.3	30.0	24.0	35.3	1.3	30.2
OV-72B	81.3	64.0	46.7	41.3	30.7	32.7	46.0	50.0	44.5
GPT-4o	63.3	54.7	44.7	40.7	55.0	42.0	54.0	32.2	46.2
Gem-1.5P	74.3	60.1	49.7	36.7	52.3	43.5	52.3	50.3	49.3
Human	-	-	91.5	92.9	92.6	92.9	89.9	91.5	100

Table 2. Results on VELOCITI using the **StrictVLE** evaluation strategy. The tests are abbreviated as Ctrl (Control), AgRand (Agent Random), AgBind (Agent Binding), AgCref (Agent Coreference), ActAdv (Action Adversarial), ActMan (Action Manner), ActBind (Action Binding), EvChr (Event Chronology). Avg reports the average accuracy on the 7 tests of VELOCITI. All models show a large gap to human performance.

achieves 49.3%, far from human performance at 93.0%. VELOCITI is a challenging benchmark and exposes lack of reasoning in both open and closed Video-LLMs.

Agent understanding tests include the Agent Random Test (AgRand), Agent Binding Test (AgBind), and Agent Coreference Test (AgCref). Broadly, they evaluate a model’s ability to understand the *doer* of the actions in the videos. Compared to AgRand, models show worse performance on AgBind and AgCref. For example, the best performing OV-72B, achieves 63.7%, 45.4%, and 38.6% accuracy respectively. AgRand requires verifying the *presence* of the agent, AgBind requires disambiguating between people present in the video and *binding* the correct person with the event description, and AgCref needs resolving identity across *multiple events*. This proves the difficulty of in-video negation. We also note that OV-7B performs worse than OV-72B on complex tests (AgCref, 8.0% vs. 38.6%) indicating that multi-level reasoning is slightly better with larger LLMs.

Action understanding tests include the Action Adversarial Test (ActAdv), Action Manner Test (ActMan), and Action Binding Test (ActBind). These evaluate the model’s understanding of actions and/or its modifiers. We observe that OV-72B scores worse on action tests (35.8% average over the 3 tests) as compared to agent tests (49.2%), while GPT-4o achieves a balanced performance of 50.3% and 46.7% respectively. While ActAdv is easier than ActBind for most models, OV shows inverted results. Further, subtle variations in actions are not captured by most models and ActMan is a challenging test with OV-72B at 29.3% and Gem-

1.5P posting the highest accuracy of 43.5%.

Multi-event understanding tests. As AgBind and ActBind adopt in-video negation, they require some level of multi-event reasoning, but are ignored in this discussion. Instead, we focus on Agent Coreference Test (AgCref) and Event Chronology Test (EvChr) as they have multiple events in both captions. Time and event order are critical to video comprehension. However, Video-LLMs are still poor at the EvChr test that requires establishing the relative order of two events. Apart from OV-72B (46.5%) and Gem-1.5P (subset, 50.3%), all models are comparable to or worse than random. This is likely as all entities mentioned in both captions are present in the video. AgCref fairs slightly better, with more models showing performance better than random: QVL-72B (35.3%), OV-72B (38.6%), Gem-1.5P (36.7%), GPT-4o (40.7%). However, it is concerning that the smaller OV-7B model collapses on these tests (AgCref 8.0%, EvChr 30.5%). Both tests highlight challenges of reasoning across multiple events in Video-LLMs.

Negation strategies. Finally, we observe that tests adopting in-video negation and requiring associations are harder than text-inspired negation. For GPT-4o that achieves balanced accuracy on agent and action understanding, we observe a 5.5% drop in performance (54.9% AgRand+ActAdv to 49.4% AgBind+ActBind)³. Solving tests with in-video negation requires reasoning as it is insufficient to only check presence of entities (all entities from both captions appear in the video). Models need to go beyond detecting the agent and action, and learn to associate them correctly.

Qualitative analysis. Example predictions of OV-72B for each test are in Fig. 7, Fig. 8, Fig. 9.

5.2. Analyzing Entailment Scores for StrictVLE

We analyze whether a model is better at classifying C^+ or C^- in Sec. 5.2. The first number in each table cell corresponds to the accuracy of positive captions, while the second number is the accuracy of negative captions when the positive caption was correct. We see an interesting trend. As the model size increases, the positive caption accuracy decreases (85.9% $\rightarrow$ 80.4%) and negative caption accuracy increases (38.1% $\rightarrow$ 53.6%). This holds for both variants: OV-7B to OV-72B and QVL-7B to QVL-72B (although on a subset). Small models are eager to say ‘Yes’ for both captions, while larger models reason better. A similar trend is seen on the control tests for the 7B and 72B models. However, negative caption accuracies are far higher, confirming why control tests are easier compared to our benchmark.

Somewhat unexpectedly, GPT-4o only achieves 64.5% accuracy on positive captions. But among them, it gets the highest negative caption accuracy of 72.3%. This hesita-

³The chosen tests provide a head-to-head comparison of text-inspired vs. in-video negation with the same positive caption as seen in Fig. 2.

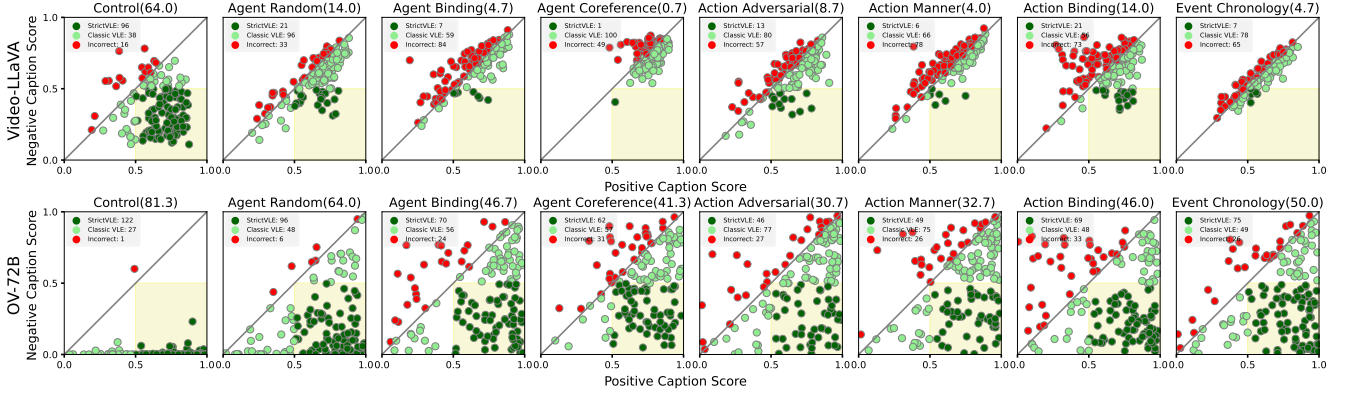


Figure 3. Scatter plot of entailment scores $e(V, C^+)$ (x-axis) and $e(V, C^-)$ (y-axis) for all tests in VELOCITI subset. We visualize the scores for Video-LLaVA (top) and OV-72B (bottom). ClassicVLE calls a sample correct in the region below the diagonal (light green). Instead, StrictVLE requires the dots to lie in the yellow bottom-right quadrant (dark green). Finally, samples whose points are above the diagonal are wrong for both VLE metrics (red). While recent models have improved, older models concentrate near the diagonal and in the top-right ‘Yes’ quadrant for both captions. The legend includes the actual number of points (please zoom in). Figure is best seen in color.

Model	Control	Average
OV-7B	83.0 / 98.3	85.9 / 38.1
OV-72B	79.8 / 99.3	80.4 / 53.6
QVL-7B	92.4 / 91.5	93.9 / 17.1
VELOCITI Subset		
QVL-72B	84.0 / 98.4	85.2 / 36.4
Gem-1.5F	93.9 / 97.8	95.8 / 25.4
GPT-4o	63.0 / 100.	64.5 / 72.3
Gem-1.5P	75.0 / 99.1	74.0 / 66.2

Table 3. StrictVLE Analysis. We study a model’s failure modes via positive and negative caption accuracy. Each cell shows the fraction of: (i) correctly classified positive captions; and (ii) correctly classified negative captions among samples whose positive captions are correct.

tion to say ‘Yes’ hurts GPT-4o on the control tests as well and even though negative caption accuracy is perfect, it gets many positive captions wrong. Similar analysis of each test (Tab. 7) shows that harder tests tend to have lower negative caption accuracies.

5.3. Evaluation with ClassicVLE

While we recommend StrictVLE, we present results on the ClassicVLE setup (Tab. 4) for completeness. First, we present a language only baseline that evaluates if C^+ is more plausible than C^- . VERA [30] scores 58.3% (close to random 50%), confirming that language biases are insufficient to solve the tests. Next, we evaluate CLIP-style models [36, 49, 50] that mean-pool video frames and observe a small improvement (ViFi-C 61.2%). New Video-LLMs such as QVL and OV (OV-72B: 83.6%) show good improvement over older ones (e.g. P-LLaVA: 59.5%). However, this score is worse than 99.4% on the easy control tests. Even with a relaxed metric, OV-72B gets every sixth sample wrong. The trends for agent and action understanding are similar: AgRand > AgBind > AgCref, and QVL and OV perform better on agent than action understanding.

To further analyze entailment scores, we present scatter plots on the benchmark subset in Fig. 3. While OV-72B is

Model	Ctrl	Ag Rand	Ag Bind	Ag Cref	Act Adv	Act Man	Act Bind	Ev Chr	Avg
Random	50.0	50.0	50.0	50.0	50.0	50.0	50.0	50.0	50.0
VERA	50.9	58.4	53.4	63.7	67.6	58.3	53.3	53.4	58.3
SigLIP	95.3	79.0	54.4	50.4	66.4	55.0	54.0	48.2	58.2
ViFi-C	93.7	82.8	58.9	56.3	63.2	60.3	59.0	48.1	61.2
Neg-C	93.4	83.5	55.3	50.4	61.6	61.1	52.4	50.1	59.2
P-LLaVA	90.7	74.6	48.9	63.7	71.0	57.0	51.8	49.8	59.5
V-LLaVA	89.7	75.1	49.6	64.3	61.6	48.5	52.0	53.8	57.8
Owl-Con	90.8	73.2	49.9	48.1	72.4	61.8	52.7	42.5	57.2
QVL-7B	97.7	93.0	74.6	63.7	75.3	76.2	70.0	63.5	73.8
OV-7B	98.6	94.4	78.8	69.0	79.7	76.9	74.2	84.0	79.6
OV-72B	99.4	95.8	83.3	80.5	84.2	81.2	78.4	81.9	83.6

Table 4. Evaluation with **ClassicVLE**. Random accuracy is 50%. Beyond Video-LLMs, we report results for a plausibility-evaluation model (VERA) and contrastive models (SigLIP: ViT-SO400M-14-SigLIP-384 [50], ViFi-C: VIFICLIP-B16 [36], and Neg-C: NegCLIP-B32 [49]). The performance of contrastive models and older Video-LLMs is close to random. However, recent models (e.g. OV) produce better relative entailment scores, even if they generate incorrect ‘Yes’/‘No’ responses.

clearly better than Video-LLaVA, it has too many points in the top-right quadrant indicating a bias to say ‘Yes’ to both captions. For Video-LLaVA, it is concerning that scores are close to the diagonal (i.e. both C^+ and C^- get similar entailment scores). In fact, these plots motivate us to propose StrictVLE and reveal problems hidden by ClassicVLE. We visualize such plots for all models in Fig. 4.

5.4. MC Evaluations and Choice Bias

In this setup, we provide the video and both captions to the Video-LLM and ask it to pick the correct description (A or B, prompts in Fig. 15, Fig. 16). In Tab. 5, we report accuracy of the model where C^+ is option A or option B.

Model	Control				Benchmark Average			
	A	B	Bias	A $\wedge$ B	A	B	Bias	A $\wedge$ B
Random	50.0	50.0	-	25.0	50.0	50.0	-	25.0
QVL-7B	94.9	98.5	(+3.6)	94.6	38.9	88.0	(+49.1)	38.5
OV-7B	96.0	99.6	(+3.6)	95.9	28.7	96.5	(+67.8)	28.7
OV-72B	99.2	99.4	(+0.2)	99.0	78.0	88.3	(+10.3)	76.0
VELOCITI Subset								
QVL-72B	100.	100.	(+0.0)	100.	73.1	75.7	(+2.6)	65.3
OV-72B	100.	100.	(+0.0)	100.	77.1	87.8	(+10.7)	75.0
Gem-1.5F	100.	99.3	(-0.7)	99.3	85.1	73.2	(-11.9)	67.7
GPT-4o	100.	100.	(+0.0)	100.	83.9	74.8	(-9.1)	68.6

Table 5. Multi-choice (MC) evaluation results. Along with video, we provide the model both captions as A and B and ask it to pick the better aligned one. Column headers A (or B) refer to the accuracy when A (or B) is the positive caption. Bias is B minus A and should be close to 0. A $\wedge$ B involves evaluating the model twice, once with correct caption as A and again as B. A sample is deemed correct when it picks the correct choice in both cases. While a model’s decision should be unaffected by the order in which choices are presented, we see a considerable bias.

We see that small 7B models have a strong choice bias and pick option B more than A (49.1% QVL-7B or 67.8% OV-7B). While this reduces in larger models (10.7% OV-72B), it is still high. Even closed models exhibit this behavior with Gem-1.5F preferring option A over B (11.9%) and GPT-4o preferring option A over B (9.1%). Interestingly, this bias becomes a major issue when the tests are challenging and is negligible in the control tests that are easier.

While one could report accuracy by running the model twice, once with option A as C^+ and again with B as C^+ (referred as A $\wedge$ B), this is tedious and the number of evaluations increases as a factorial of the number of choices. If we compare StrictVLE with the MC evaluation’s A $\wedge$ B score (both apply $\wedge$ on binary decisions and have random chance at 25%), we observe that MC is much easier (OV-72B: 76.0%) than StrictVLE (43.1%, Tab. 2). This may be attributed to the MC setup, where a model processes both captions at once and only needs to pick the more likely option; in contrast with the StrictVLE setup that requires independent evaluation of each caption. Even though MC is a popular evaluation setup for many benchmarks (see Tab. 1), the choice bias of Video-LLMs makes results difficult to interpret. For all these reasons, StrictVLE is preferred.

5.5. Validating Benchmark Properties

We highlight some additional properties of our benchmark.

Evaluating blind models. Tab. 6A compares Qwen2-LLM (Q LLM) and OV-72B without the video inputs (OV Blind) against the default OV-72B model (here, OV ). We see a dramatic drop (43.1% to 3.7% Q LLM and 8.1% OV Blind). Solving tests in VELOCITI requires visual understanding.

Evaluating with a single-frame. Tab. 6B reports results on

Model	Ctrl	Ag Rand	Ag Bind	Ag Cref	Act Adv	Act Man	Act Bind	Ev Chr	Avg
A. Comparing against Blind Models									
OV 	79.3	63.7	45.4	38.6	33.1	29.3	45.1	46.5	43.1
Q LLM	2.2	2.5	2.2	5.3	4.1	1.3	2.7	7.9	3.7
OV Blind	6.0	9.3	6.7	12.7	10.3	3.3	6.9	7.4	8.1
B. Impact of Single Frame Input or Model									
OV-7B	81.6	56.7	32.9	8.0	29.7	30.6	36.4	30.5	32.1
1 Frame	39.6	31.6	22.3	18.1	15.5	13.4	22.1	10.3	19.0
OV-7B-SI	78.9	35.7	15.6	2.9	27.6	21.4	22.0	8.8	19.1

Table 6. VELOCITI benchmark validation. **Part A.** We confirm that the model requires visual inputs. The base LLM Qwen2-72B (Q LLM) or OV-72B without providing the video (OV Blind) perform poorly compared to OV-72B provided with video frames (OV ). **Part B.** We also confirm that providing multiple video frames is necessary. When OV-7B is provided a single frame chosen randomly (1 Frame) or when the video is fed to a model trained only on single images (LLaVA-OneVision-SingleImage, OV-7B-SI), the performance dips compared to showing the video at 1fps (our default strategy, OV-7B).

OV-7B models. The first row is the default setup (Tab. 2) and is compared against: (i) OV-7B with a single frame input (1 Frame) chosen at random from the sampled 1fps frames. (ii) The OneVision team [23] first train an image-only model and extend it to multiple images and videos. We evaluate their single image checkpoint while providing video inputs (OV-7B-SI). The performance drops from 32.1% to about 19.0% in both cases. This confirms that VELOCITI requires video inputs and video models.

CoT and FPS ablations are in Appendix A.3 and A.4.

6. Conclusion

We introduced VELOCITI, a benchmark to evaluate the compositional capabilities of Video-LLMs by disentangling and assessing the comprehension of agents, actions, and their associations across multiple events. We improved over the classic Video-Language Entailment setup that relies on relative scoring by proposing StrictVLE that requires models to answer ‘Yes’ for the positive caption *and* ‘No’ for the negative caption. All evaluated models, open and closed, performed poorly with a large gap to human performance. Our experiments showed that action understanding is harder than agent understanding, and solving tests with in-video negation is harder than text-inspired ones. We also analyzed limitations of ClassicVLE and the choice bias in multiple-choice evaluations. Overall, our work established that compositional reasoning on short videos is still unsolved and remains challenging for Video-LLMs.

Acknowledgment. MT and ViG thank Adobe Research for supporting this project and travel. MT thanks SERB SRG/2023/002544 for compute server and Google Cloud credits. We thank volunteers for human evaluations.

References

- [1] Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altenschmidt, Sam Altman, Shyamal Anadkat, et al. GPT-4 Technical Report. *arXiv preprint arXiv:2303.08774*, 2023. 5
- [2] Meta AI. Llama3. <https://llama.meta.com/llama3/>, 2024. 3
- [3] Piyush Bagad, Makarand Tapaswi, and Cees GM Snoek. Test of Time: Instilling Video-Language Models With a Sense of Time. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2023. 2, 3, 5
- [4] Hritik Bansal, Yonatan Bitton, Idan Szepktor, Kai-Wei Chang, and Aditya Grover. VideoCon: Robust Video-Language Alignment via Contrast Captions. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2023. 2, 3, 5
- [5] Lucas Beyer, Andreas Steiner, André Susano Pinto, Alexander Kolesnikov, Xiao Wang, Daniel Salz, Maxim Neumann, Ibrahim Alabdulmohsin, Michael Tschannen, Emanuele Bugliarello, et al. PaliGemma: A versatile 3B VLM for transfer. *arXiv preprint arXiv:2407.07726*, 2024. 1
- [6] Shyamal Buch, Cristóbal Eyzaguirre, Adrien Gaidon, Jiajun Wu, Li Fei-Fei, and Juan Carlos Niebles. Revisiting the “Video” in Video-Language Understanding. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2022. 2, 3
- [7] Mu Cai, Reuben Tan, Jianrui Zhang, Bocheng Zou, Kai Zhang, Feng Yao, Fangrui Zhu, Jing Gu, Yiwu Zhong, Yuzhang Shang, Yao Dou, Jaden Park, Jianfeng Gao, Yong Jae Lee, and Jianwei Yang. TemporalBench: Towards Fine-grained Temporal Understanding for Multimodal Video Models. *arXiv preprint arXiv:2410.10818*, 2024. 1
- [8] Daniel Cores, Michael Dorkenwald, Manuel Mucientes, Cees GM Snoek, and Yuki M Asano. TVBench: Redesigning Video-Language Evaluation. *arXiv preprint arXiv:2410.07752*, 2024. 2, 3
- [9] Wenliang Dai, Junnan Li, Dongxu Li, Anthony Tiong, Junqi Zhao, Weisheng Wang, Boyang Li, Pascale Fung, and Steven Hoi. InstructBLIP: Towards General-purpose Vision-Language Models with Instruction Tuning. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2023. 1
- [10] Google Deepmind. Gemini 1.5: Unlocking multimodal understanding across millions of tokens of context. *arXiv preprint arXiv: 2403.05530*, 2024. 2, 5
- [11] Matt Deitke, Christopher Clark, Sangho Lee, Rohun Tripathi, Yue Yang, Jae Sung Park, Mohammadreza Salehi, Niklas Muennighoff, Kyle Lo, Luca Soldaini, et al. Molmo and PixMo: Open weights and open data for state-of-the-art multimodal models. *arXiv preprint arXiv:2409.17146*, 2024. 1
- [12] Fartash Faghri, David J Fleet, Jamie Ryan Kiros, and Sanja Fidler. Vse++: Improving visual-semantic embeddings with hard negatives. In *British Machine Vision Conference (BMVC)*, 2018. 2
- [13] Xinyu Fang, Kangrui Mao, Haodong Duan, Xiangyu Zhao, Yining Li, Dahua Lin, and Kai Chen. MMBench-Video: A Long-Form Multi-Shot Benchmark for Holistic Video Understanding. In *Advances in Neural Information Processing Systems (NeurIPS): Track on Datasets and Benchmarks*, 2024. 1, 3, 24
- [14] Andrea Frome, Greg S Corrado, Jon Shlens, Samy Bengio, Jeff Dean, Marc’Aurelio Ranzato, and Tomas Mikolov. Devise: A deep visual-semantic embedding model. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2013. 2
- [15] Chaoyou Fu, Yuhao Dai, Yongdong Luo, Lei Li, Shuhuai Ren, Renrui Zhang, Zihan Wang, Chenyu Zhou, Yunhang Shen, Mengdan Zhang, Peixian Chen, Yanwei Li, Shaohui Lin, Sirui Zhao, Ke Li, Tong Xu, Xianwu Zheng, Enhong Chen, Rongrong Ji, and Xing Sun. Video-MME: The First-Ever Comprehensive Evaluation of Multi-modal LLMs in Video Analysis. *arXiv preprint*, 2024. 1, 2, 3, 24
- [16] Madeleine Grunze-McLaughlin, Ranjay Krishna, and Maneesh Agrawala. AGQA: A Benchmark for Compositional Spatio-Temporal Reasoning. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2021. 2, 3
- [17] Cheng-Yu Hsieh, Jieyu Zhang, Zixian Ma, Aniruddha Kembhavi, and Ranjay Krishna. SugarCreme: Fixing Hackable Benchmarks for Vision-Language Compositionality. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2024. 2
- [18] Daniel Jurafsky and James H. Martin. *Speech and Language Processing: An Introduction to Natural Language Processing, Computational Linguistics, and Speech Recognition with Language Models*, 2024. 4
- [19] Muhammad Uzair Khattak, Muhammad Ferjad Naeem, Jameel Hassan, Muzammal Naseer, Federico Tomba, Fahad Shahbaz Khan, and Salman Khan. How Good is my Video LMM? Complex Video Reasoning and Robustness Evaluation Suite for Video-LMMs. *arXiv preprint arXiv:2405.03690*, 2024. 2, 3
- [20] Martha Lewis, Nihal Nayak, Peilin Yu, Jack Merullo, Qinan Yu, Stephen Bach, and Ellie Pavlick. Does CLIP Bind Concepts? Probing Compositionality in Large Image Models. In *European Chapter of the Association for Computational Linguistics (EACL)*, 2024. 2
- [21] Bohao Li, Yuying Ge, Yixiao Ge, Guangzhi Wang, Rui Wang, Ruimao Zhang, and Ying Shan. SEED-Bench-2: Benchmarking Multimodal Large Language Models. *arXiv preprint arXiv:2311.17092*, 2023. 2, 3
- [22] Baiqi Li, Zhiqiu Lin, Wenxuan Peng, Jean de Dieu Nyandwi, Daniel Jiang, Zixian Ma, Simran Khanuja, Ranjay Krishna, Graham Neubig, and Deva Ramanan. NaturalBench: Evaluating Vision-Language Models on Natural Adversarial Samples. *arXiv preprint arXiv:2410.14669*, 2024. 3
- [23] Bo Li, Yuanhan Zhang, Dong Guo, Renrui Zhang, Feng Li, Hao Zhang, Kaichen Zhang, Yanwei Li, Ziwei Liu, and Chunyuan Li. LLaVA-OneVision: Easy Visual Task Transfer. *arXiv preprint arXiv:2408.03326*, 2024. 1, 2, 5, 8
- [24] 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 *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2024. 2, 3

- [25] Shicheng Li, Lei Li, Shuhuai Ren, Yuanxin Liu, Yi Liu, Rundong Gao, Xu Sun, and Lu Hou. VITATECS: A Diagnostic Dataset for Temporal Concept Understanding of Video-Language Models. In *European Conference on Computer Vision (ECCV)*, 2024. 2, 3, 5, 26
- [26] Yifan Li, Yifan Du, Kun Zhou, Jinpeng Wang, Wayne Xin Zhao, and Ji-Rong Wen. Evaluating Object Hallucination in Large Vision-Language Models. In *Empirical Methods in Natural Language Processing (EMNLP)*, 2023. 2
- [27] Yunxin Li, Xinyu Chen, Baotian Hu, Longyue Wang, Haoyuan Shi, and Min Zhang. VideoVista: A Versatile Benchmark for Video Understanding and Reasoning. *arXiv preprint arXiv:2406.11303*, 2024. 3
- [28] Bin Lin, Bin Zhu, Yang Ye, Munan Ning, Peng Jin, and Li Yuan. Video-LLaVA: Learning United Visual Representation by Alignment Before Projection. *arXiv preprint arXiv:2311.10122*, 2023. 1, 5
- [29] Zhiqiu Lin, Deepak Pathak, Baiqi Li, Jiayao Li, Xide Xia, Graham Neubig, Pengchuan Zhang, and Deva Ramanan. Evaluating Text-to-Visual Generation with Image-to-Text Generation. In *European Conference on Computer Vision (ECCV)*, 2024. 3
- [30] Jiacheng Liu, Wenya Wang, Dianzhuo Wang, Noah A Smith, Yejin Choi, and Hannaneh Hajishirzi. Vera: A General-Purpose Plausibility Estimation Model for Commonsense Statements. In *Empirical Methods in Natural Language Processing (EMNLP)*, 2023. 7
- [31] 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? In *Findings of the Association for Computational Linguistics*, 2024. 3
- [32] Zixian Ma, Jerry Hong, Mustafa Omer Gul, Mona Gandhi, Irena Gao, and Ranjay Krishna. CREPE: Can Vision-Language Foundation Models Reason Compositionally? In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2023. 2
- [33] Jae Sung Park, Sheng Shen, Ali Farhadi, Trevor Darrell, Yejin Choi, and Anna Rohrbach. Exposing the Limits of Video-Text Models through Contrast Sets. In *North American Chapter of Association of Computational Linguistics: Human Language Technologies (NAACL-HLT)*, 2022. 2, 3
- [34] 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. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2024. 3
- [35] Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, et al. Learning Transferable Visual Models From Natural Language Supervision. In *International Conference on Machine Learning (ICML)*, 2021. 2
- [36] Hanoona Rasheed, Muhammad Uzair Khattak, Muhammad Maaz, Salman Khan, and Fahad Shahbaz Khan. Fine-Tuned CLIP Models Are Efficient Video Learners. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2023. 7
- [37] Ruchit Rawal, Khalid Saifullah, Ronen Basri, David Jacobs, Gowthami Somepalli, and Tom Goldstein. CinePile: A Long Video Question Answering Dataset and Benchmark. *arXiv preprint arXiv:2405.08813*, 2024. 1, 3
- [38] Arijit Ray, Filip Radenovic, Abhimanyu Dubey, Bryan Plummer, Ranjay Krishna, and Kate Saenko. Cola: A Benchmark for Compositional Text-to-image Retrieval. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2024. 2
- [39] Arka Sadhu, Tanmay Gupta, Mark Yatskar, Ram Nevatia, and Aniruddha Kembhavi. Visual Semantic Role Labeling for Video Understanding. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2021. 2, 3, 4, 23, 26
- [40] Sanyal, Soumya and Xiao, Tianyi and Liu, Jiacheng and Wang, Wenya and Ren, Xiang. Are Machines Better at Complex Reasoning? Unveiling Human-Machine Inference Gaps in Entailment Verification. In *Findings of the Association for Computational Linguistics*, 2024. 2, 5
- [41] Gunnar A Sigurdsson, Gül Varol, Xiaolong Wang, Ali Farhadi, Ivan Laptev, and Abhinav Gupta. Hollywood in Homes: Crowdsourcing Data Collection for Activity Understanding. In *European Conference on Computer Vision (ECCV)*, 2016. 2
- [42] Tristan Thrush, Ryan Jiang, Max Bartolo, Amanpreet Singh, Adina Williams, Douwe Kiela, and Candace Ross. Winoground: Probing Vision and Language Models for Visio-Linguistic Compositionality. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2022. 2, 5
- [43] Maxim Tkachenko, Mikhail Malyuk, Andrey Holmanyuk, and Nikolai Liubimov. Label Studio: Data labeling software, 2020-2022. Open source software available from <https://github.com/heartexlabs/label-studio>. 15
- [44] Shengbang Tong, Zhuang Liu, Yuexiang Zhai, Yi Ma, Yann LeCun, and Saining Xie. Eyes Wide Shut? Exploring the Visual Shortcomings of Multimodal LLMs. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2024. 2
- [45] Peng Wang, Shuai Bai, Sinan Tan, Shijie Wang, Zhihao Fan, Jinze Bai, Keqin Chen, Xuejing Liu, Jialin Wang, Wenbin Ge, Yang Fan, Kai Dang, Mengfei Du, Xuancheng Ren, Rui Men, Dayiheng Liu, Chang Zhou, Jingren Zhou, and Junyang Lin. Qwen2-VL: Enhancing Vision-Language Model's Perception of the World at Any Resolution. *arXiv preprint arXiv:2409.12191*, 2024. 1, 5
- [46] Bo Wu, Shoubin Yu, Zhenfang Chen, Joshua B. Tenenbaum, and Chuang Gan. STAR: A Benchmark for Situated Reasoning in Real-World Videos. In *Advances in Neural Information Processing Systems (NeurIPS): Track on Datasets and Benchmarks*, 2021. 2, 3
- [47] Ning Xie, Farley Lai, Derek Doran, and Asim Kadav. Visual Entailment Task for Visually-grounded Language Learning. In *Visually Grounded Interaction and Language (ViGIL) Workshop at NeurIPS*, 2018. 2, 3, 5
- [48] Lin Xu, Yilin Zhao, Daquan Zhou, Zhijie Lin, See Kiong Ng, and Jiashi Feng. PLLaVA: Parameter-free LLaVA Extension

- from Images to Videos for Video Dense Captioning. *arXiv preprint arXiv:2404.16994*, 2024. [1](#), [5](#)
- [49] Mert Yuksekgonul, Federico Bianchi, Pratyusha Kalluri, Dan Jurafsky, and James Zou. When and Why Vision-Language Models Behave like Bags-Of-Words, and What to Do About It? In *International Conference on Learning Representations (ICLR)*, 2022. [2](#), [7](#)
 - [50] Xiaohua Zhai, Basil Mustafa, Alexander Kolesnikov, and Lucas Beyer. Sigmoid Loss for Language Image Pre-Training . In *International Conference on Computer Vision (ICCV)*, 2023. [7](#)
 - [51] Jianrui Zhang, Cai Mu, and Yong Jae Lee. Vinoground: Scrutinizing LMMs over Dense Temporal Reasoning with Short Videos. *arXiv*, 2024. [2](#), [3](#)
 - [52] Tiancheng Zhao, Tianqi Zhang, Mingwei Zhu, Haozhan Shen, Kyusong Lee, Xiaopeng Lu, and Jianwei Yin. VL-Checklist: Evaluating Pre-trained Vision-Language Models with Objects, Attributes and Relations. *arXiv preprint arXiv:2207.00221*, 2022. [2](#)

Supplementary Material

In this supplementary material, we discuss

1. Additional results and analysis, both quantitative and qualitative (Appendix A);
2. Benchmark creation, quality control process, and some statistics (Appendix B);
3. Model prompts used in both setups: entailment and multiple-choice (Appendix C); and
4. Limitations (Appendix D).

A. Additional Results

In Appendix A.1, we present scatter plots of entailment scores for all models across all tests, expanding Fig. 3 from the main paper. Next, we present the positive and negative entailment scores that are used in StrictVLE (expanding the analysis in Sec. 5.2) in Appendix A.2. We experiment with Chain-of-Thought prompting in Appendix A.3 and present ablations for number of sampled frames in Appendix A.4. The multiple-choice (MC) evaluation results are discussed in Appendix A.5 and the human evaluation setup in Appendix A.6. Finally, we share some qualitative results of LLaVA-OneVision-72B on our benchmark in Appendix A.7.

A.1. Scatter plot of entailment scores

To analyze entailment scores, we present scatter plots for all models on the benchmark subset (150 samples) in Fig. 4. The ideal scenario is when all samples lie in the bottom-right quadrant (points in dark green, quadrant in light yellow), which indicates that the model confidently entails the correct caption while rejecting the negative caption, leading to a 100% StrictVLE accuracy. However, in practice, we observe two undesirable cases: (i) the points are concentrated in the top-right quadrant, indicating a strong bias towards responding ‘Yes’ regardless of whether the caption is aligned or misaligned; and (ii) the points are clustered around the diagonal, indicating that the model exhibits similar confidence levels when saying ‘Yes’ to both the positive and negative captions. Major takeaways are highlighted below:

- P-LLaVA has most of its points concentrated in the top-right quadrant, indicating a strong bias towards responding ‘Yes’ regardless of whether the caption is positive or negative, which also explains its near 0% StrictVLE accuracy.
- Owl-Con and Video-LLaVA are strongly clumped near the diagonal in the top-right quadrant (except for the Control Test): indicating that they tend to respond ‘Yes’ and have similar entailment scores for both the positive and negative captions. Owl-Con appears to be worse than Video-LLaVA with more points in the top-right quadrant.
- Between LLaVA-OneVision-7B (OV-7B) and LLaVA-OneVision-7B-Si (OV-7B-SI), we see that the points in OV-7B-SI are more clustered near the diagonal while LLaVA-OneVision-7B is more diffused except for AgCref. This is expected as it is hard for a model trained on single images to distinguish between the positive and the negative caption and nearly impossible for the EvChr and AgCref. In contrast, both models perform well on the Control Test since the replacements come from a totally different or random video, making it easier for the models to classify with sufficient confidence.
- For Qwen2-VL-7B (QVL-7B) except for Control Test, the points for all the other tests are concentrated in the top-right corner while additionally being clustered near the diagonal for the EvChr. QVL-7B performs worse than OV-7B even though both models are trained using the same base Qwen2 7B parameter LLM.
- Finally, on LLaVA-OneVision-72B, we see that many points are below the diagonal and would score correct on ClassicVLE. However, roughly half of them (on average) are in the bottom right quadrant indicating difficulty of the best model to predict ‘Yes’ for the positive caption and ‘No’ for the negative caption respectively.

A.2. Analyzing Entailment Scores for StrictVLE

Continuing from findings of Sec. 5.2 in the main paper, we analyze whether a model finds it easier to classify C^+ or C^- in Tab. 7 for all tests. Each cell in the table reports two numbers: the first is the accuracy of positive captions, and the second is the accuracy of negative captions when the positive caption is correct.

An interesting observation (as also noted in the main paper) is that as model size increases, the positive caption accuracy decreases while the negative caption accuracy improves. This holds for both variants: OV-7B to OV-72B and QVL-7B to QVL-72B, and indicates that small models are eager to say ‘Yes’ for both captions, while larger models reason better.

Although Qwen2-VL (QVL) models achieve higher accuracy for positive captions than LLaVA-OneVision (OV) models, the negative caption accuracy is better for OV models. This indicates that QVL models are biased to say ‘Yes’ regardless of the captions, whereas OV models reason better and are less inclined to respond ‘Yes’. For QVL, specifically for the EvChr

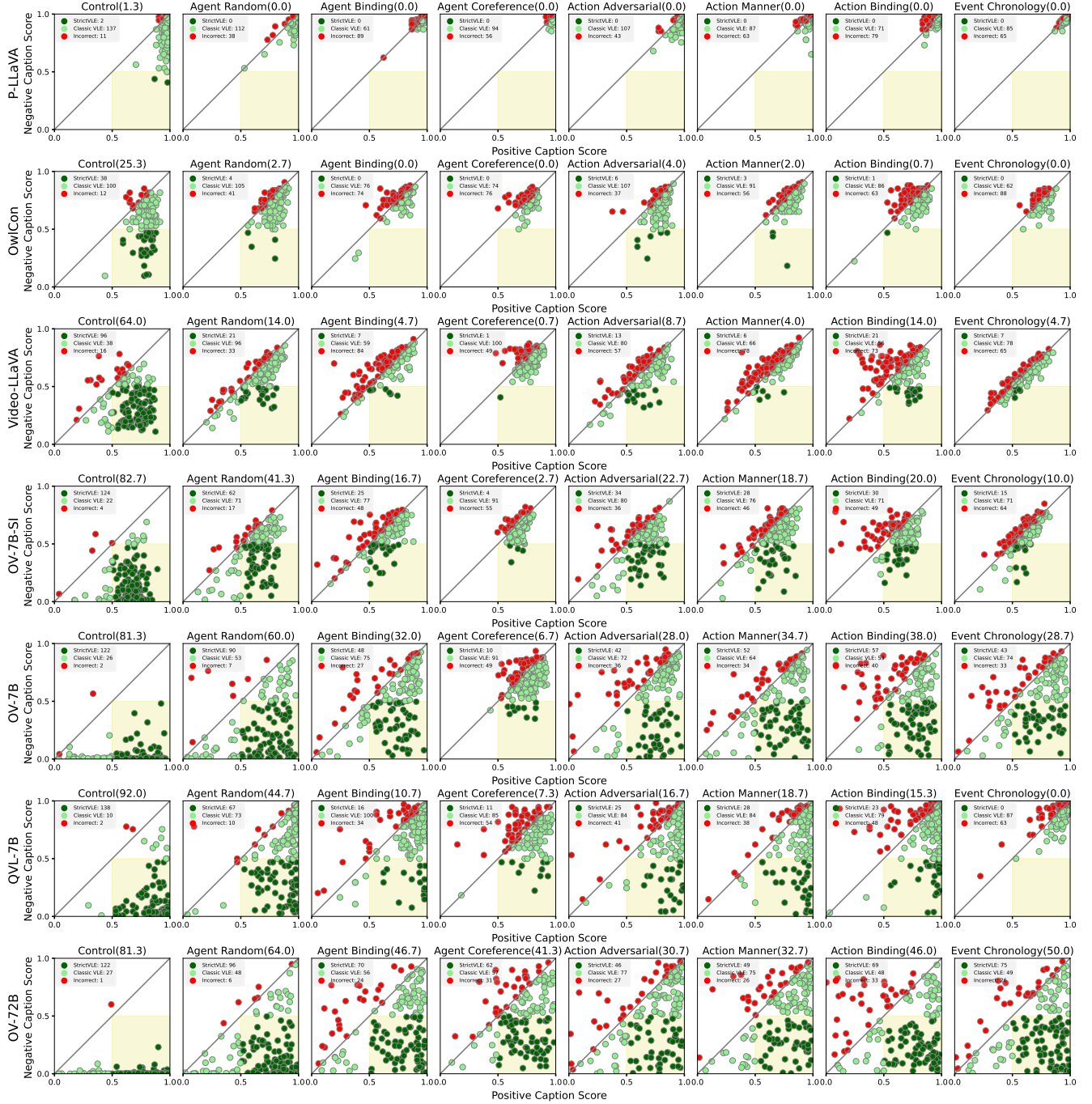


Figure 4. Scatter plot of entailment scores $e(V, C^+)$ (x-axis) and $e(V, C^-)$ (y-axis) for all tests in VELOCITY. We visualize the scores for several models indicated in the left margin. From top to bottom: P-LLaVA, OwlCon, Video-LLaVA, OV-7B-SI, OV-7B, QVL-7B, and OV-72B. ClassicVLE calls a sample correct in the region below the diagonal (light green). Instead, StrictVLE requires the dots to lie in the yellow bottom-right quadrant (dark green). Finally, samples whose points are above the diagonal are wrong for both VLE metrics (red). The legend includes the actual number of points (please zoom in). This figure is best seen in color.

test, the positive caption accuracy is very high, but the negative caption accuracy is extremely low, indicating that the QVL models are very poor at the temporal order reasoning.

While GPT-4o achieves comparatively lower accuracy on positive captions across all tests, it consistently achieves the highest accuracy for negative captions, except in the EvChr test, where OV-72B performs best. Another surprising observation

Model	Ctrl	Ag Rand	Ag Bind	Ag Cref	Act Adv	Act Man	Act Bind	Ev Chrono	Avg
OV-7B	83.0 / 98.3	84.2 / 67.3	82.0 / 40.1	97.1 / 8.2	86.3 / 34.4	80.6 / 37.9	81.9 / 44.5	89.0 / 34.2	85.9 / 38.1
OV-72B	79.8 / 99.3	81.6 / 78.1	79.5 / 57.1	87.0 / 44.4	80.6 / 41.1	77.9 / 37.5	78.2 / 57.6	78.3 / 59.4	80.4 / 53.6
QVL-7B	92.4 / 91.5	92.8 / 42.1	90.7 / 14.9	97.6 / 6.6	94.3 / 18.9	93.0 / 18.8	91.0 / 18.1	98.0 / 0.4	93.9 / 17.1
VELOCITI Subset									
QVL-72B	84.0 / 98.4	85.3 / 65.6	84.0 / 34.9	77.3 / 45.7	84.0 / 35.7	86.7 / 27.7	80.7 / 43.8	98.7 / 1.4	85.2 / 36.4
OV-72B	81.3 / 100.	79.3 / 80.7	80.7 / 57.9	86.7 / 47.7	79.3 / 38.7	82.0 / 39.8	74.7 / 61.6	80.7 / 62.0	80.5 / 55.5
Gem-1.5F	93.9 / 97.8	93.3 / 60.4	93.9 / 25.4	100. / 4.7	93.3 / 35.3	95.3 / 22.7	95.3 / 26.2	99.3 / 2.8	95.8 / 25.4
Gem-1.5P	75.0 / 99.1	72.3 / 83.2	75.5 / 65.8	71.3 / 51.4	72.5 / 72.2	79.6 / 54.7	75.5 / 65.8	71.4 / 70.5	74.0 / 66.2
GPT-4o	63.3 / 100.0	58.0 / 94.3	64.0 / 69.8	62.0 / 65.6	65.8 / 83.7	65.3 / 64.3	64.7 / 83.5	71.8 / 44.9	64.5 / 72.3

Table 7. StrictVLE Analysis for various models on all tests in VELOCITI. Each cell of the table has two numbers. The first is the fraction of correctly classified positive captions. The second is the fraction of correctly classified negative captions, among samples whose positive caption is classified correctly. Refer to Appendix A.2 for a description.

Model	Ag Rand	Ag Bind	Ag Cref	Act Adv	Act Man	Act Bind	Ev Chr	Avg
OV-72B								
w/o CoT	64.0	46.7	41.3	30.7	32.7	46.0	50.0	44.5
CoT	40.0	28.0	19.3	26.7	28.0	32.0	30.0	29.1
Gemini-1.5 Pro								
w/o CoT	60.1	49.7	36.7	52.3	43.5	52.3	50.3	49.3
CoT	46.6	32.8	45.5	46.2	46.8	46.4	29.2	41.9

Table 8. Average score on VELOCITI subset: without and with CoT

is that Gemini-1.5-Flash (Gem-1.5F), despite achieving the best accuracy for positive captions, performs worse than all other models for negative captions. This suggests that Gemini-1.5-Flash may also be responding with ‘Yes’ too often. Additionally, both Gemini-1.5-Flash and Qwen2-VL-72B exhibit very low accuracy for negative captions in the AgCref and EvChr tests.

Finally, in Sec. 5.1 of the main paper, we highlight that $AgRand > AgBind > AgCref$ – this trend is clearly observed in the negative caption accuracies presented in the Table, and explains the poor performance of some models on Agent Coreference Test (single-digit accuracies on negative captions).

A.3. Impact of Chain-of-Thought prompting

We experimented with Chain-of-Thought (CoT) prompting for Gemini-1.5-Pro and OV-72B (prompt in Fig. 5). As shown in Tab. 8, the performance reduced in both cases indicating that models are unable to reason in a step-by-step manner for such statements.

A.4. Impact of Increased Frame Rate

We explore increasing the video sampling rate to observe if more visual information aids the model to solve the tasks in VELOCITI to a greater extent. For this, we sample frames at 8 fps, amounting to 80 frames for a 10 s video. From Tab. 9 we observe that the smaller model (OV-7B) benefits with more frames resulting in improvement across most tests, an average of +3.5%. Interestingly, its larger counterpart (OV-72B) performs worse with significant drops on action tests, ActAdv and ActMan, (9-11%). This may be due to the large context size that the model is not trained for. Both models perform better on EvChr task.

A.5. Multiple-Choice (MC) Evaluation: Results on each test

In the MC setup, we provide the video along with both captions to the Video-LLM and ask it to pick the correct one (A or B). Results on the control and average over the benchmark were discussed in Sec. 5.4 of the main paper.

Now, we report results across all the tests in Tab. 10. For both OV and QVL models, we see that the smaller variants have a higher choice bias and tend to prefer option B. While this bias reduces in the larger variants, it is still high. Also, as

System Prompt

You are an AI assistant specializing in analyzing movie clips to verify captions using a Chain of Thought (CoT) approach. Given a movie clip and a corresponding caption, your task is to determine whether the caption accurately describes the events in the clip.

A caption is considered accurate (“Yes”) if **all applicable** of the following criteria are met:

1. ****Actor/Doer****: The person or entity performing the action is correctly identified.
2. ****Attributes****: The characteristics of the actors/doers and the action itself are accurately described (*e.g.* clothing color, size, speed).
3. ****Instruments/Objects****: Any tools, objects, or instruments used in the action are correctly identified.
4. ****Receiver/Patient****: The target or recipient of the action is correctly identified.
5. ****Relationships****: The relationships between the entities involved (*e.g.*, “standing next to”, “holding”) are accurately depicted.
6. ****Manner****: The way in which the action is performed (*e.g.*, “quickly”, “slowly”, “angrily”) is accurately described.
7. ****Location****: The setting or location of the scene is correctly identified.
8. ****Clarity****: There is sufficient visual information in the clip to confidently assess the correctness of the caption.
9. ****Event Order****: If the caption suggests a specific order of events, then the video should have events happening in the suggested order.

If any of the above criteria cannot be verified due to a lack of visuals, the caption should not be considered accurate.

Note that the caption is designed to represent a part of the video clip and may not explain all the events in the clip.

Follow these steps:

1. ****Analysis****: Carefully examine the provided movie clip.
2. ****Reasoning****: Analyze the caption in relation to the clip. Break down the caption into smaller parts and determine if each part meets the accuracy criteria listed above. Detail your reasoning process within ‘<thinking>’ tags.
3. ****Evaluation****: Based on your reasoning, evaluate the overall accuracy of the caption. If there is insufficient information in the clip to definitively confirm or deny the caption based on one or more criteria, explain what information is missing within ‘<reflection>’ tags.
4. ****Conclusion****: Provide a clear “Yes” or “No” answer within ‘<output>’ tags.

Use the following format:

<thinking>

[Detailed step-by-step reasoning, referencing the accuracy criteria. This is your internal thought process.]

</thinking>

<reflection>

[Reflections on your reasoning, including any uncertainties or missing information and which criteria could not be verified. If the caption cannot be definitively verified, explain why.]

</reflection>

<output>

[Yes or No]

</output>

Evaluate the following caption for the accompanying movie clip: {caption}

Figure 5. CoT evaluation prompt.

expected, the accuracy of $A \wedge B$ improves for larger variants. We observe that harder tests (*e.g.* AgBind vs. AgRand) tend to have a higher bias. Among all the tests, the EvChr test has the highest bias and the lowest accuracy across all the models.

Both Gemini-1.5-Flash and GPT-4o show considerable bias. Interestingly, GPT-4o seems to prefer option A, while Gemini-1.5F prefers option B.

A.6. Human Evaluation

Human evaluations were conducted in a standardized manner to establish human performance in the various tasks presented in VELOCITY. The evaluations included 3 volunteers who were assigned the subset (150 samples for each of the 7 tests). This amounts to a total of 2,100 video-caption pairs (7 tests $\times$ 150 samples $\times$ 2 captions). We use the Label Studio [43] annotation platform for this task. To ensure fair evaluations, humans are first shown a set of instructions to ensure consistency across

Model	Ag Rand	Ag Bind	Ag Cref	Act Adv	Act Man	Act Bind	Ev Chr	Avg
OV-7B								
1fps	56.7	32.9	8.0	29.7	30.6	36.4	30.5	32.1
8fps	59.3	34.7	6.0	34.7	38.3	33.3	42.0	35.6
OV-72B								
1fps	64.7	46.0	36.7	42.0	40.7	46.0	46.0	46.0
8fps	63.7	45.4	38.6	33.1	29.3	45.1	46.5	43.1

Table 9. Higher frame rate sampling results.

Model	AgRand		AgBind		AgCref		ActAdv		ActMan		ActBind		EvChr	
	Bias	A^B	Bias	A^B	Bias	A^B	Bias	A^B	Bias	A^B	Bias	A^B	Bias	A^B
QVL-7B	24.2	74.1	42.2	40.8	37.5	33.6	49.6	42.9	51.5	41.5	40.0	36.1	98.5	0.7
OV-7B	41.6	58.1	81.6	17.1	59.9	26.0	71.3	27.6	68.0	30.1	70.0	24.2	81.9	15.9
OV-72B	3.1	94.8	10.9	72.9	8.0	69.0	8.9	79.7	11.1	77.5	14.8	62.5	15.1	75.9
VELOCITI Subset														
QVL-72B	6.0	88.7	3.3	64.7	2.7	60.0	6.7	68.0	2.0	74.0	8.6	54.7	-11.3	47.3
OV-72B	2.7	95.3	11.4	75.3	11.3	67.3	9.3	76.7	7.3	77.3	16.0	61.3	16.7	72.0
Gem-1.5F	-2.8	94.4	-12.6	73.4	8.0	61.3	-12.9	72.8	-14.3	66.0	0.7	64.8	-49.6	41.4
GPT-4o	4.1	92.5	4.7	79.3	3.4	60.8	-10.1	77.0	-7.0	74.1	2.0	70.7	-60.8	25.7

Table 10. MC evaluation results on all tests. Along with the video, we provide the model with both captions A and B and ask it to pick the better-aligned one. Bias is the accuracy difference between B and A options and should be close to 0. A^B involves evaluating the model twice, once with the correct caption as A and again as B. A sample is deemed correct when it picks the correct choice in both cases. While a model’s decision should be unaffected by the order in which choices are presented, a considerable bias is observed.

participants. Next, we randomize and present non-overlapping video-caption pairs. An example of the annotation dashboard is shown in Fig. 6.

A.7. Qualitative Analysis

We present examples from the OV-72B model on our benchmark for three following cases: (i) Samples satisfying the StrictVLE criteria ($e(V, C^+) > 0.5 \wedge e(V, C^-) < 0.5$) are shown in Fig. 7; (ii) Samples *only* satisfying the ClassicVLE condition ($e(V, C^+) > e(V, C^-)$), but failing on the StrictVLE condition are in Fig. 8. (iii) Finally, samples classified incorrectly according to ClassicVLE ($e(V, C^+) < e(V, C^-)$), are presented in Fig. 9. Note these are also incorrect for StrictVLE. In each case, we show 10 frames from the video, the positive and negative captions, and the corresponding entailment scores. The test name is indicated in the bottom left.

Instructions

These instructions can be opened anytime by clicking 'i' on the bottom left of the panel.

You are given a video and a caption for each task.

Please watch the 10s video and select 'Yes' if the given video entails the caption, otherwise select 'No'

- The caption should provide an accurate description of the events in the video.
- The caption should correctly identify the entities (*humans, animals, objects, etc.*) and the relationships (*actions*) between them.

Note

- Ignore any spelling/grammatical errors, if any.
- You may watch the video multiple times, if needed.

Video 1



1 of 241

00:00:00 00:10:00

Here is a caption that describes the video

The woman with black hair is inserting a tongue suppressor into the girl with brown hair's mouth.

Does the given video entail the caption?

☐ Yes^[1] ☐ No^[2]

Binary-Choice, Single Select Option

Revisit Annotation Instructions

Skip Submit

Figure 6. Human Evaluation Dashboard. Instructions and interface for human evaluation for the entailment task.

	
Control	
$e(V, C^+) = 0.944$	$e(V, C^-) = 0.000$
	
Agent Random	
$e(V, C^+) = 0.841$	$e(V, C^-) = 0.468$
	
Agent Binding	
$e(V, C^+) = 0.603$	$e(V, C^-) = 0.055$
	
Agent Coreference	
$e(V, C^+) = 0.640$	$e(V, C^-) = 0.268$
	
Action Adversarial	
$e(V, C^+) = 0.885$	$e(V, C^-) = 0.075$
	
Action Manner	
$e(V, C^+) = 0.743$	$e(V, C^-) = 0.124$
	
Action Binding	
$e(V, C^+) = 0.987$	$e(V, C^-) = 0.060$
	
Event Chronology	
$e(V, C^+) = 0.798$	$e(V, C^-) = 0.480$

Figure 7. VELOCITI samples where OV-72B classifies the sample correctly based on the StrictVLE criteria. In the [Agent Binding](#) example, the scene visualizes a man and a woman talking on the phone while the man drives, C^- changes the entity of the driver. The model is confidently able to identify that it is the man who is driving and not the woman, as the positive caption scores (0.603) much above the negative caption (0.055) while satisfying the StrictVLE criteria. Similarly, in [Agent Coreference](#), the scene describes two women - a woman in blue who's sitting and puts on her headphones as she begins to write, while the woman in white looks at her and eventually walks away. The C^- interchanges the roles of these two women, and the model correctly scores the positive caption (0.640) higher than the negative caption (0.268).

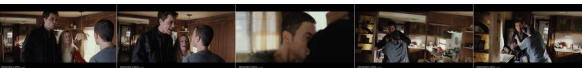
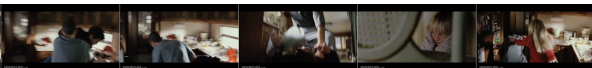
 Positive Caption		 Negative Caption	
In an apartment, a woman bends down to speak to a duck		At the side of the house, a woman knocks down items on the counter with her hand	
Control		$e(V, C^+) = 0.734$ $e(V, C^-) = 0.527$	
 Positive Caption		 Negative Caption	
Agent Random		$e(V, C^+) = 0.637$ $e(V, C^-) = 0.519$	
 Positive Caption		 Negative Caption	
Agent Binding		$e(V, C^+) = 0.911$ $e(V, C^-) = 0.5$	
 Positive Caption		 Negative Caption	
Agent Coreference		$e(V, C^+) = 0.861$ $e(V, C^-) = 0.661$	
 Positive Caption		 Negative Caption	
Action Adversarial		$e(V, C^+) = 0.679$ $e(V, C^-) = 0.569$	
 Positive Caption		 Negative Caption	
Action Manner		$e(V, C^+) = 0.969$ $e(V, C^-) = 0.721$	
 Positive Caption		 Negative Caption	
Action Binding		$e(V, C^+) = 0.870$ $e(V, C^-) = 0.531$	
 Positive Caption		 Negative Caption	
Event Chronology		$e(V, C^+) = 0.845$ $e(V, C^-) = 0.644$	
First, In a bar, a ghost wearing glasses is chugging alcohol from a glass. Then, In a bar, the ghost in glasses is removing his hat from his head		First, In a bar, the ghost in glasses is removing his hat from his head. Then, In a bar, a ghost wearing glasses is chugging alcohol from a glass	

Figure 8. VELOCITI samples where OV-72B classifies the sample correctly based on the ClassicVLE criteria, but not on StrictVLE. In the **Action Binding** example, a man in a black plaid jacket is cowering. The negative caption (C^-) changes the action from “cowering” to “talking aggressively.” Although the model assigns a high entailment score of 0.870 to the positive caption (C^+), it also assigns a relatively high score of 0.531 to the negative caption (C^-). While this satisfies the ClassicVLE criterion, it fails to meet the StrictVLE criterion.

	
Positive Caption	Negative Caption
Outside a building, a woman is reaching down with her hand to grab a bag	In a warehouse, the boy is approaching the man laying on the bed with caution
Control	
$e(V, C^+) = 0.144$	$e(V, C^-) = 0.162$
	
Positive Caption	Negative Caption
Outside a mansion, a woman opens the door to see who's there	Outside the mansion, the bald man is opening the door to see who's there
Agent Random	
$e(V, C^+) = 0.091$	$e(V, C^-) = 0.668$
	
Positive Caption	Negative Caption
Inside a restaurant, the woman in a black shirt exits the kitchen	Inside the restaurant, the man in a brown jacket is exiting the kitchen
Agent Binding	
$e(V, C^+) = 0.746$	$e(V, C^-) = 0.779$
	
Positive Caption	Negative Caption
The person who is aiding the man in a yellow shirt is also the one who is removing his stethoscope from his ears	The person who is aiding the man in a yellow shirt is also the one who is covering his mouth, looking distraught
Agent Coreference	
$e(V, C^+) = 0.503$	$e(V, C^-) = 0.554$
	
Positive Caption	Negative Caption
In a bedroom, the dog is walking forward to get out from under the covers	In a bedroom, the dog is settling into the covers
Action Adversarial	
$e(V, C^+) = 0.787$	$e(V, C^-) = 0.933$
	
Positive Caption	Negative Caption
At the back seat of a car, the blonde woman is picking up a bag with her left hand	The blonde woman picks up the bag from the back seat of the car with both hands
Action Manner	
$e(V, C^+) = 0.647$	$e(V, C^-) = 0.888$
	
Positive Caption	Negative Caption
In a backyard, the man wearing white pants is watching the man wearing a grey tank top as he holds food in his hand	In a backyard, the man wearing white pants is chopping wood
Action Binding	
$e(V, C^+) = 0.074$	$e(V, C^-) = 0.166$
	
Positive Caption	Negative Caption
First, A man in a green jacket is throwing an upper cut punch to a policeman's face. Then, At a police station, a policeman is scowling at a man in a green jacket	First, At a police station, a policeman is scowling at a man in a green jacket. Then, A man in a green jacket is throwing an upper cut punch to a policeman's face
Event Chronology	
$e(V, C^+) = 0.081$	$e(V, C^-) = 0.134$

Figure 9. VELOCITI samples classified incorrectly even for ClassicVLE. In **Agent Random**, the scene describes a woman opening the door for a man and hugging him. C^- replaces the person opening the door with a random person (a bald man), and the model makes a mistake - scoring the negative caption (0.668) considerably more than the positive caption (0.091). **Action Manner** has a video of two women driving into the scene where a blonde woman picks up a bag from the backseat using her left arm. The C^- modifies how the bag is picked up - with both hands, which is clearly incorrect. However, the model makes a mistake and prefers the negative caption (0.888) over the positive caption (0.647).

B. Benchmark Creation and Details

In this section, we provide details about our benchmark. In particular, we share all prompts used for creating positive captions and various tests (Appendix B.1, Appendix B.2), share our process on creating a benchmark subset for evaluating closed models (Appendix B.3), provide benchmark statistics (Appendix B.4), discuss the strategy used to manually verify and clean all the tests (Appendix B.5), and finally provide some compute and runtime details that are required to evaluate on our benchmark (Appendix B.6).

B.1. Prompt for Converting SRL Dictionary to a Positive Caption

The prompt for generating the positive caption given an SRL dictionary is shown in Fig. 10. This refers to the discussion from Sec. 3.1 in the main paper. We use a two-stage strategy that first inserts all elements of the SRL dictionary in a sentence and then refines it for proper grammatical structure.

B.2. Prompts for Creating Test Samples

The prompt above (Fig. 10) helps create the positive caption for multiple tests. Specifically, Agent Random Test, Agent Binding Test, Action Adversarial Test, Action Manner Test, and Action Binding Test, all use the above strategy, while Agent Coreference Test and Event Chronology Test adopt templates that are filled in with the complete (or partial) positive captions.

The negative prompts for Agent Random Test, Agent Binding Test, and Action Binding Test are also created in the same

System Prompt

Using the provided dictionary containing verb and argument-role pairs in the style of PropBank, follow these steps to generate two captions

Naive Caption: Generate a caption that faithfully reflects all details from the dictionary without adding or omitting any information. Ensure that every argument detail is accurately included in the Naive Caption.

Fluent Caption: If the Naive Caption is already fluent and naturally phrased, directly copy it to the Fluent Caption. If necessary, refine the Naive Caption for improved language fluency while strictly maintaining all original details and arguments from the dictionary.

Please proceed with generating the Naive Caption first, ensuring it remains comprehensive and accurate based on the provided dictionary entries. Then, if adjustments are needed to enhance fluency, refine the Naive Caption into the Fluent Caption while ensuring that no details are overlooked or omitted.

Few Shot Example 1

```
{'Verb': 'walk (walk)',  
'Arg0 (walker)': 'man in suit',  
'ArgM (direction)': 'into room',  
'ArgM (manner)': 'slowly',  
'Scene of the Event': 'Warehouse'}
```

Naive Caption: In a warehouse, a man in suit is walking slowly into the room.

Fluent Caption: In a warehouse, a man in suit is walking slowly into the room.

Few Shot example 2

```
{'Verb': 'burn (cause to be on fire)',  
'Arg0 (thing burning)': 'Wreckage',  
'ArgM (location)': 'Wreckage'}
```

Naive Caption: The wreckage is burning on the wreckage.

Fluent Caption: The wreckage is burning.

Figure 10. Prompt to generate the positive caption given an SRL dictionary.

System Prompt

Your objective is to generate a contradiction caption using the provided PropBank style “input dictionary” and the ‘Verb’ labelled as ‘source’ based on a specific “misalignment scenario” called “verb misalignment”. In this scenario, you should suggest an alternative contradictory value for the “source” and label it as “target”.

Key Requirements

1. “naive caption + verb misalignment”: should be plausible and could theoretically occur in real life.
2. The “fluent caption + verb misalignment”: If the “naive caption + verb misalignment” is already fluent and naturally phrased, directly copy it to the “fluent caption + verb misalignment”. If necessary, refine the “naive caption + verb misalignment” for improved language fluency while strictly maintaining all original details and arguments from the dictionary

Guidelines

1. The “target” should introduce a contradiction when compared to “source”, without being a mere negation.
2. The “naive caption + verb misalignment” should be clearly distinguishable from the scene described by the “input dictionary” and should be visually distinguishable.
3. Your replacements should be creative yet reasonable.
4. If adjustments are needed to enhance fluency, refine the “naive caption + verb misalignment” into the “fluent caption + verb misalignment” while ensuring that no details are overlooked or omitted.

Few Shot Example 1

```
{'Verb': 'speak (speak)',  
'Arg0 (talker)': 'a man with dark hair',  
'Arg2 (hearer)': 'old man'  
'ArgM (manner)': 'greeting him',  
'Scene of the Event': 'warehouse'}
```

Target: Ignore

Naive Caption: On the front porch, a man with dark hair is ignoring an old man, greeting him.

Fluent Caption: On the front porch, a man with dark hair is ignoring an old man.

Few Shot Example 2

```
{'Verb': 'open (open)',  
'Arg0 (opener)': 'woman with long hair',  
'Arg1 (thing opening)': 'the front door',  
'ArgM (manner)': 'slowly',  
'Scene of the Event': 'inside a house'}
```

Target: Close

Naive Caption: Inside a house, a woman with long hair is closing the front door slowly.

Fluent Caption: Inside a house, a woman with long hair is closing the front door slowly.

Figure 11. Prompt to generate the negative caption for Action Adversarial Test.

manner as above by first replacing the specific Verb or Arg0 in the dictionary followed by strategy above.

Finally, the prompt for generating the Action Adversarial Test negative caption is shown in Fig. 11 and for Action Manner Test negative captions in Fig. 12. Both involve generating a target replacement that seems reasonable followed by converting the SRL dictionary into a caption.

B.3. Subset Creation

We created a subset of VELOCITI with 150 samples in each test. The subset was curated through random tries such that the StrictVLE performance of the OV-72B model was comparable to the full set, allowing for fair comparisons.

System Prompt

Your objective is to generate a contradiction caption using the provided PropBank style “input dictionary” and the ‘ArgM (manner)’ labeled as ‘source’ based on a specific “misalignment scenario” called “manner misalignment”. In this scenario, you should suggest an alternative contradictory value for the “source” and label it as “target”

Key Requirements

1. “naive caption + manner misalignment”: should be plausible and could theoretically occur in real life.
2. The “fluent caption + manner misalignment”: If the “naive caption + manner misalignment” is already fluent and naturally phrased, directly copy it to the “fluent caption + manner misalignment”. If necessary, refine the “naive caption + manner misalignment” for improved language fluency while strictly maintaining all original details and arguments from the dictionary.

Guidelines

1. The “target” should introduce a contradiction when compared to “source”, without being a mere negation.
2. The “naive caption + manner misalignment” should be clearly distinguishable from the scene described by the “input dictionary.”
3. Your replacements should be creative yet reasonable.
4. If adjustments are needed to enhance fluency, refine the “naive caption + manner misalignment” into the “fluent caption + manner misalignment” while ensuring that no details are overlooked or omitted

Few Shot Example 1

```
{'Verb': 'look (vision)',  
'Arg0 (looker)': 'a man wearing all black',  
'Arg1 (thing looked at or for or on)': 'a building'  
'ArgM (direction)': 'infront of him',  
'ArgM (manner)': 'breathing heavily',  
'Scene of the Event': 'warehouse'}
```

Target: *Whistling*

Naive Caption: Outside, a man wearing all black is looking in front of him at a building while whistling.

Fluent Caption: Outside, a man wearing all black is looking at a building in front of him while whistling.

Few Shot Example 2

```
{'Verb': 'burn (cause to be on fire)',  
'Arg0 (agent, entity causing something to be suspended)': 'climbing ropes',  
'Arg1 (thing suspended)': 'woman in pink shirt',  
'Arg2 (suspended from)': 'climbing ropes',  
'ArgM (location)': 'on the face of the rocks',  
'ArgM (manner)': 'precariously'}
```

Target: *Securely*

Naive Caption: climbing ropes are hanging the woman in a pink shirt securely on the face of the rocks.

Fluent Caption: The woman in a pink shirt is hanging on the face of the rocks from the climbing ropes securely.

Figure 12. Prompt to generate the negative caption for Action Manner Test.

B.4. Benchmark Statistics

We present some statistics highlighting the diversity and nuance in the VELOCITI benchmark. Since this benchmark is a subset of VidSitu [39], we observe similar trends as presented in their work.

Videos in our benchmark are complex as there are multiple agents performing various actions. Actions in VELOCITI are fine-grained. We analyze the set using Gemini-1.5-Pro which broadly categorizes actions into 6 groups: physical action and movement, communication and expression, manipulation and physical interaction, perception and mental activity, physiological actions, and general activities and states. In general, models struggle slightly more with physiological actions (performance $\sim 10\%$ lower) as compared to the average. Some verbs from these categories are shown in Fig. 13, note that the size of the word here does *not* correspond to its frequency in the dataset.



Figure 13. Word-cloud of some actions in VELOCITI in different action categories as suggested by Gemini-1.5 Pro. Word **size does not correspond to frequency** and is assigned randomly for visualization.

Fig. 14a shows that around 87% of the videos contain 4 or more unique verbs, and Fig. 14b shows that about 85% of videos contain 2 or more unique agents (people performing actions). We evaluate binding by leveraging the fact that one agent can perform multiple actions in the video, and the richness of the SRL annotations ensure that these events are described adequately. In Fig. 14c, we observe that over 70% of the events contain 4 or more SRLs (*e.g.* agent, patient, manner, *etc.*), indicating the detail-oriented nature of the annotations. Finally, Fig. 14d shows that over 72% of agents occur twice or more in their corresponding video annotation. These agents would likely be performing two different actions, and we utilize this to create two references to the same agent in tests such as Agent Coreference Test.

B.5. Quality Control

To ensure that the data generated from the automated pipelines discussed earlier are correct, we filtered the data samples manually, following specific guidelines discussed in this section. The final count of the data samples is reported in Tab. 11.

SRL dictionary to caption. The instructions and the interface for evaluating caption quality is described in Fig. 17. For each sample, three choices were provided: positive if the caption is correct, negative if the caption is wrong, and neutral if the caption cannot be negative but contains some ambiguity due to which it could not be considered positive. Out of the 380 samples that were manually verified, 356 were marked as positive, 21 were neutral, and 3 were negative. The number of positive and neutral samples was high (99.2%).

All tests. For each sample of all tests, we perform a meticulous cleanup. The instructions and the interface are presented in Fig. 18. For each video, the green bar contains a positive caption, and the red bar contains a negative caption. Unlike human evaluations, the positive and the negative captions are known while filtering. Only the samples for which both positive and negative captions are deemed appropriate are retained.

B.6. Runtime and Compute Details

While benchmarks on long videos are interesting [13, 15], VELOCITI proposes important challenges that every Video-LLM needs to solve. The short 10 s videos enable fast evaluation and make the benchmark accessible: running OV-7B on all tests (except the Control Test) takes about 2.6 hours on a single RTX 4090 GPU (24 GB).

C. Model Evaluation Prompts

We present the prompts used for all open Video-LLMs, Gemini-1.5-Flash, and GPT-4o. The entailment and MC evaluation prompts for open models, such as Qwen2-VL, LLaVA-OneVision, and Gemini-1.5-Flash are provided in Fig. 15. Prompts

Test	Videos	# Samples	Subset
Ctrl	850	2635	150
AgRand	588	873	150
AgBind	615	1459	150
AgCref	183	339	150
ActAdv	355	438	150
ActMan	378	458	150
ActBind	540	1356	150
EvChr	521	1234	150

Table 11. Number of videos and samples across different tests in VELOCITI.

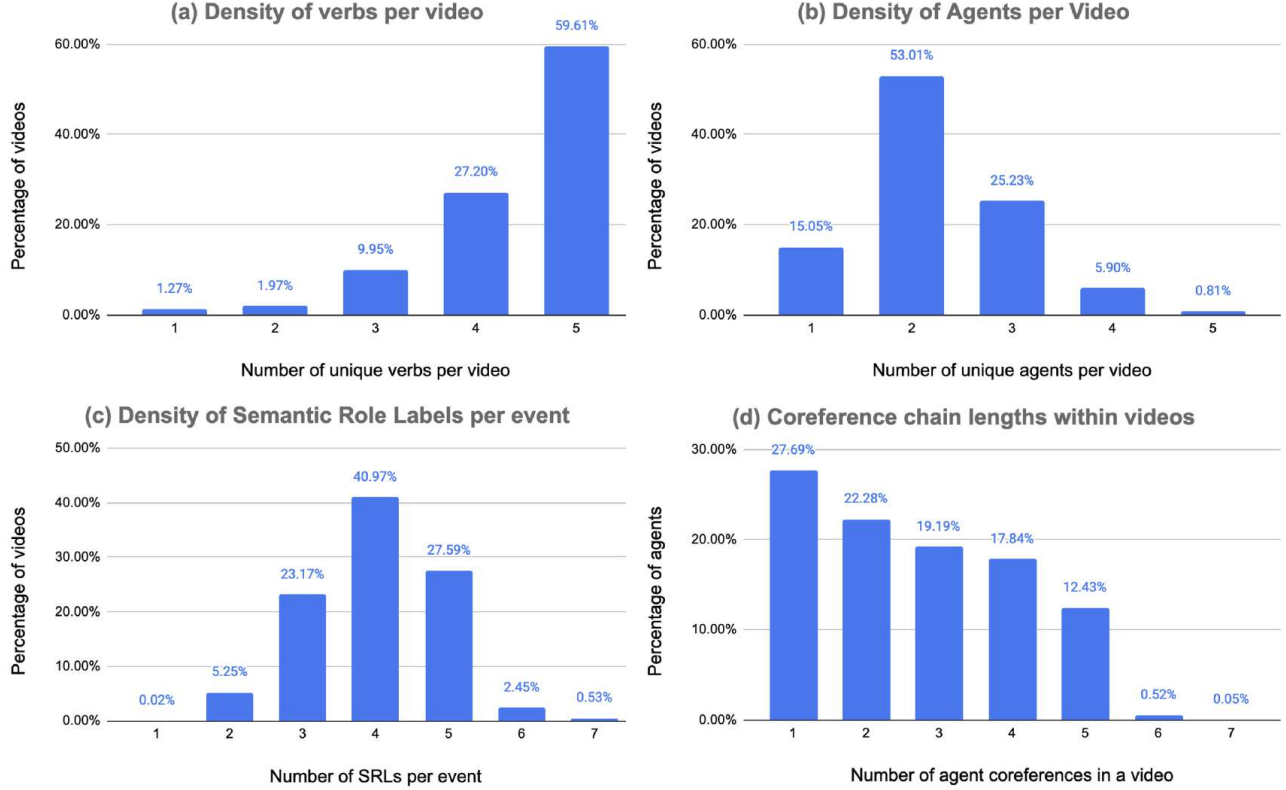


Figure 14. Statistics of various features of the VELOCITI benchmark. (a) and (b) show the distribution of verbs and agents per video, respectively. (c) shows the density of SRL annotations per event; and (d) shows the distribution of agent coreference lengths. Even with short videos, the complexity of the VidSitu annotations make the task challenging.

for GPT-4o are shown in Fig. 16. Note that GPT-4o is provided the explicit instruction of being provided frames of a video, while others are directly given a video.

Although some closed models have started optionally sharing logits, they are restricted to a limited top-K set, *e.g.* top-20 for GPT-4o. Hence, the logits for the ‘Yes’ and ‘No’ tokens may not always be included in these top-k values. To ensure the evaluation of closed models covers maximum data samples, the prompts were slightly modified to explicitly include the instruction: “Just answer with either Yes or No.”.

Entailment Prompt

Carefully watch the video and pay attention to the sequence of events, the details and actions of persons.
Here is a caption that describes the video: *Caption*
Based on your observation, does the given video entail the caption?

MC Prompt

Carefully watch the video and pay attention to the sequence of events, the details and actions of persons.
Here are two captions that describe the video.
A) *Caption₁*
B) *Caption₂*
Based on your observation, select the caption that best describes the video.
Just print either A or B.

Figure 15. Prompts for **Open Video-LLMs** and **Gemini-1.5-Flash** and **Gemini-1.5-Pro**. **Top:** Entailment evaluation prompt. **Bottom:** Multiple-choice evaluation prompt.

Entailment Prompt

You are given frames sampled sequentially from a video. Carefully watch the video frames and pay attention to the sequence of events, the details and actions of persons.

Here is a caption that describes the video: *Caption*

Based on your observation, does the given video entail the caption?

Just answer with either Yes or No.

MC Prompt

You are given frames sampled sequentially from a video. Carefully watch the video frames and pay attention to the sequence of events, the details and actions of persons.

Here are two captions that describe the video.

A) *Caption₁*

B) *Caption₂*

Based on your observation, select the caption that best describes the video.

Just print either A or B.

Figure 16. Prompts for **GPT-4o**. **Top**: Entailment evaluation prompt. **Bottom**: Multiple-choice evaluation prompt.

D. Limitations

We discuss some limitations of our work.

1. One of the shortcomings is the limited ability to scale the benchmark. VELOCITI relies on SRLs, which are obtained from careful (and costly) human annotations [39]. Further, we use LLMs to generate captions from the SRL dictionary and to create several tests (Appendix B.1, Appendix B.2). However, LLMs are prone to hallucinations, and hence, we do a round of human verification to confirm that the captions are appropriate. Thus, costly human intervention is required from SRL curation to verification of individual test samples.
2. VELOCITI is not intended as a one-stop benchmark to evaluate all abilities of Video-LLMs. Instead, it evaluates Video-LLMs for facets of compositionality, a fundamental aspect of visio-linguistic reasoning. Also, as VELOCITI is derived from VidSitu, a person-centric dataset, our benchmark focuses on people and their actions/interactions.
3. Lastly, our proposed StrictVLE metric cannot be used to evaluate contrastive models, as these models do not provide a direct ‘Yes’ probability. When the alignment score is used as a proxy to the entailment score (similar to [25]), we show that contrastive CLIP-based models do not perform well even with ClassicVLE and are therefore unlikely to be competitive at a stricter entailment.

Instructions

These instruction can be opened anytime by clicking 'i' on the bottom left of the panel

Your objective is to mark whether the provided positive caption is an “*accurate*” description of the dictionary contents.

What does an “*accurate*”, positive-caption mean?

- The generated caption must include all ideas *inferred* from the dictionary, even though it may miss some exact phrases.
- Ideally, the caption should include everything from the dictionary, but if the caption misses some value of an argument (for example, *direction*), then the caption is correct only if the missing value is *implied* from the caption.
- It needs to be grammatically correct, even though it may sound uncommon in conversational English.

Positive, Neutral, Negative

- Select *positive*, when the caption clearly meets the above requirements.
- Select *neutral*, when the caption partially meets the above requirement (not fully correct).
- Select *negative*, when the caption does NOT meet the above requirements.

Other Rules

- You are only required to look at the provided dictionary, and not the video for this task.
 - Captions should NOT be marked incorrect because of noisy annotations in the dictionary.
 - Captions should NOT be marked incorrect because of abrupt capitalization inside the sentence.
-

#107092418
101 of 101

Compositional Benchmark Human Evaluation

Name	Value
event	Ev2
video	v_bky6BtAbTU8_seg_85_95

Name	Value
.Verb	disbelieve (not believe)
Arg0 (non-believer)	girl in a gray hat
Scene of the Event	in the woods

In the woods, a girl in a gray hat disbelieves.

☐ Positive  ^[1] ☐ Neutral  ^[2] ☐ Negative  ^[3]

Submit

Info Comments History

Add a comment 

Regions Relations

Manual By Time 

Regions not added

Figure 17. Instructions and interface to verify the quality of captions generated from LLaMA-3-70B.

Instructions

These instructions can be opened anytime by clicking 'i' on the bottom left of the panel

You are given 1 video and 2 captions for each task, one **correct caption** and a **negative caption**.

Please watch the video and *verify* if the positive and the negative captions are “*logically correct*”.

What is a “*logically-correct*”, **positive caption**?

- Caption that provides a *correct* description of the event in the video.
- It should correctly identify the entities (*humans, animals, objects, etc.*) and the relationships (*action*) between them.
- Spelling/grammatical errors, if any, shall be ignored.

What is “*logically-correct*”, **negative caption**?

- Caption that provides an *incorrect* description of events in the video.

Note

- You may watch the video multiple times, if required.
- Careful and precise judgement is requirement, as point-of-difference between the **positive** and the **negative** caption, may be subtle.

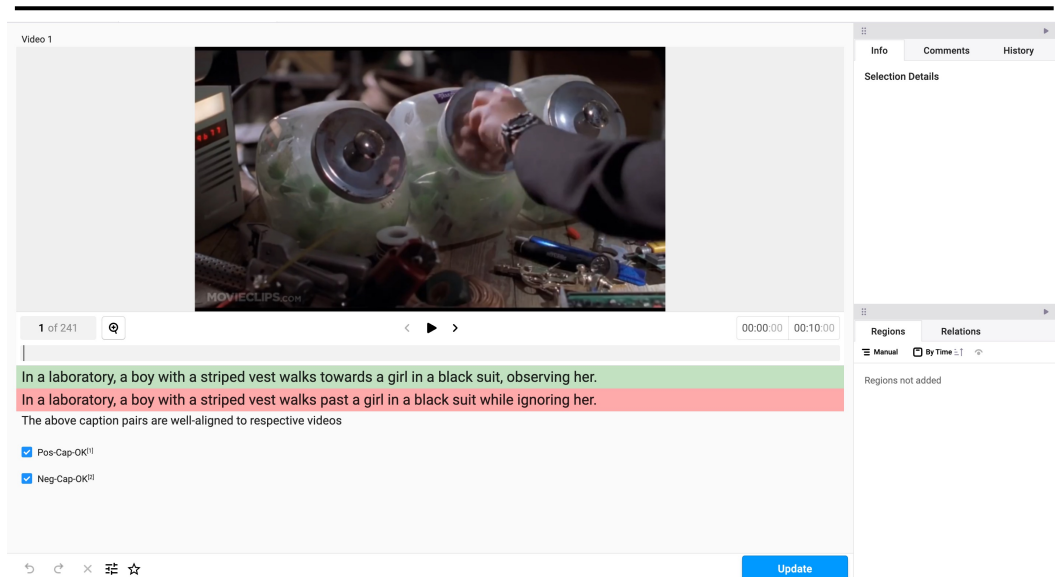


Figure 18. Data cleaning instruction for all the tests.