

VisCoP: Visual Probing for Video Domain Adaptation of Vision Language Models

Dominick Reilly¹ Manish Kumar Govind¹ Le Xue² Srijan Das¹

¹ University of North Carolina at Charlotte, Charlotte NC, USA

² Elorian AI, Palo Alto CA, USA

<https://github.com/dominickrei/VisCoP>

dreilly1@charlotte.edu

Abstract. Large Vision Language Models (VLMs) excel at general visual reasoning tasks, but their performance degrades sharply when deployed in novel domains with substantial distribution shifts compared to what was seen during pretraining. Existing approaches to adapt VLMs to novel target domains rely on finetuning standard VLM components. Depending on which components are finetuned, these approaches either limit the VLMs ability to learn domain-specific features, or lead to catastrophic forgetting of pre-existing capabilities. To address this, we introduce **V**ision **C**ontextualized **P**robing (**VisCoP**), which augments the VLM’s vision encoder with a compact set of learnable *visual probes*, enabling domain-specific features to be learned with only minimal updates to the pretrained VLM components. We evaluate VisCoP across three challenging domain adaptation scenarios: cross-view (exocentric $\rightarrow$ egocentric), cross-modal (RGB $\rightarrow$ depth), and cross-task (human understanding $\rightarrow$ robot control). Our experiments demonstrate that VisCoP consistently outperforms existing domain adaptation strategies, achieving superior performance on the target domain, while better retaining capabilities from the source domain. We will release all code, models, and evaluation protocols to facilitate future research in this direction.

Keywords: Vision-Language Models · Action Understanding

1 Introduction

Large Vision-Language Models (VLMs) [3, 35, 47, 54] have demonstrated strong performance across diverse multimodal tasks, including open-ended video question answering [29, 52] and spatial reasoning [21, 38]. These models couple pretrained vision encoders [37, 53] with Large Language Models (LLMs) [33, 36] and are trained on large-scale, web-curated image-video-text corpora containing broad but generic visual concepts [6, 31, 40, 55]. However, when deployed in domains that exhibit a pronounced *visual gap* arising from differences in viewpoint, sensing modality, or task structure, their performance degrades substantially under distribution shift. This challenge is prevalent in video domain adaptation settings such as exocentric $\rightarrow$ egocentric viewpoint, RGB $\rightarrow$ depth modality, and visual understanding $\rightarrow$ robotic control, where learning domain-specific *visual* representations is essential. Crucially, VLMs must *adapt while preserving* the

broad world knowledge acquired during pretraining: a model adapted to egocentric data, for example, should retain competence on exocentric benchmarks to maintain cross-domain generalization.

A common strategy to address such distributional shift is partial fine-tuning of a pretrained VLM on domain-specific video instruction data. Freezing the vision encoder and updating only lightweight components (e.g., the vision-language connector) preserves pretrained knowledge but restricts visual specialization. In contrast, finetuning the vision encoder enables specialized visual understanding, albeit at the cost of catastrophic forgetting of pretrained knowledge [25, 49, 51]. This trade-off is particularly severe when the dominant shift is visual. Moreover, existing approaches rely primarily on the final-layer representations of a frozen vision encoder, discarding intermediate features that encode multi-level visual structure [22]. We argue that such abstracted representations are insufficient for visually prominent domain shifts and pose the following question: *how can a VLM leverage domain-specific visual signals across the depth of a pretrained vision encoder while avoiding catastrophic forgetting?*

To this end, we introduce **Vision Contextualized Probing (VisCoP)**, a lightweight adaptation mechanism that enables pretrained VLMs to specialize to a target domain while preserving their general-purpose visual knowledge. Motivated by the progressive emergence of semantic representations across transformer depths [4, 30, 45], VisCoP performs two way probing: it first extracts domain relevant signals from intermediate layers of a frozen vision encoder, and subsequently

conditions the language model embedding space to enhance domain specific visual reasoning. This is achieved through a compact set of learnable probe tokens that interact layer wise with intermediate visual representations, forming an alternative parameterized pathway for adaptation. By leveraging multi-level abstractions across the encoder depth, VisCoP learns a parallel representation pathway that extracts domain-specific signals from intermediate encoder features while preserving pretrained representations. This enables the model to learn domain-specific visual patterns while mitigating catastrophic forgetting. In contrast, existing approaches [2, 15, 19, 24, 43] typically operate on the final-layer representations of the vision encoder or introduce prompt tokens that condition downstream reasoning without explicitly learning new visual cues *within* the vision encoder. This design provides a simple yet effective approach to video domain adaptation under prominent visual

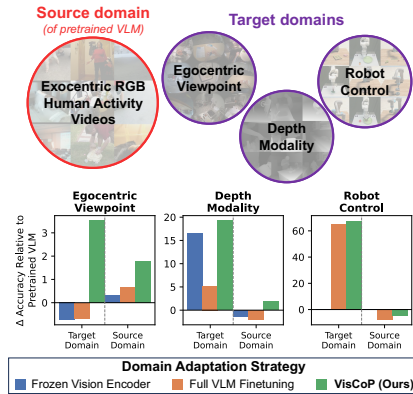


Fig. 1: Domain adaptation performance of different adaptation strategies. VisCoP achieves superior target domain performance while better retaining source domain knowledge compared to other strategies.

gaps without modifying the pretrained visual encoder parameters, which is not enabled by existing approaches. Empirically, we find that the representations learned via the VISCoP adaptation pathway enable effective cross-view, cross-modal, and cross-task adaptation of VLMs, while retaining their broad capabilities learned during pretraining. Metaphorically, the name VISCoP reflects its role as a “*traffic cop*”, directing gradient flows away from the visual encoder and towards an alternative pathway for learning domain-specific visual features, avoiding the “*crash*” (catastrophic forgetting) that would otherwise occur if gradients flowed directly through the vision encoder.

To summarize, our key contributions:

1. We propose VISCoP (**V**ision **C**ontextualized **P**robing), a novel video domain adaptation strategy for VLMs that learns domain-specific visual representations through layer-wise probing of a frozen vision encoder, enabling effective domain transfer and preventing catastrophic forgetting of multi-modal capabilities acquired during pretraining.
2. We establish a comprehensive evaluation setting for domain adaptation in VLMs, spanning three challenging target domains: cross-view (exocentric $\rightarrow$ egocentric), cross-modality (RGB $\rightarrow$ depth), and cross-task (action understanding $\rightarrow$ robotic control), along with standardized metrics to evaluate performance.
3. Our experiments demonstrate that VLMs trained with VISCoP outperform alternative domain adaptation strategies across diverse target domains, while retaining more knowledge of the source domain illustrated in Figure 1.

2 Related Works

Domain adaptation in vision-language encoders. Domain adaptation of contrastively trained vision-language encoders, such as CLIP [37, 39], is typically achieved through prompt tuning or adapter-based approaches. Both strategies aim to learn domain-specific features while keeping the pretrained vision and text encoders frozen. To accomplish this, prompt tuning approaches [50, 56, 57] introduce learnable prompt vectors as additional input to the text encoder, steering the model toward target domain. Adapter-based approaches [12, 49] insert lightweight trainable modules directly into the encoder space, thus updating their pretrained representations. Head2Toe [10] similarly exploits intermediate representations of a frozen backbone, but passively selects static activations to train a linear classifier per target task, with no focus on source-domain retention. In contrast to these approaches, VISCoP addresses the setting of domain adaptation in generative VLMs, enabling them to learn domain-specific features without requiring updates to the pretrained encoder representations.

Domain adaptation in VLMs. Domain adaptation in VLMs has largely been achieved through data-centric strategies rather than through architectural changes [7]. Existing approaches typically leverage automated pipelines [34, 41] or closed-source VLMs [5, 22] to curate visual-instruction pairs from existing datasets in the target domain. Their adaptation strategy usually follows a multi-stage training scheme similar to LLaVA [28], where different VLM components

are selectively trained at each stage. However, the choice of trainable components creates a trade-off between extracting domain-specific features and retaining pretrained knowledge. Training only lightweight connectors retains pretrained knowledge but limits domain-specific visual understanding, while training the vision encoder enables specialized visual understanding at the cost of catastrophic forgetting. VisCOP avoids this trade-off through the introduction of visual probes that extract domain-specific features from a frozen vision encoder, enabling adaptation without disrupting the pretrained visual representations.

Visual probing vs. visual compression and prompt tuning. Several approaches employ learnable tokens to bridge vision and language modalities [15, 43, 59] through architectures leveraging the Q-Former, Perceiver Resampler, or prompt tuning mechanisms. Q-Former [24] leverages learnable queries that cross-attend to representations from the final layer of the vision encoder, aggregating visual information into a reduced set of tokens for computational efficiency. Perceiver Resampler [2] operates similarly, aiming to compress the visual representations into a fixed number of learnable tokens. Prompt tuning methods [19, 27, 56] introduce learnable tokens to steer downstream reasoning, but operate solely on fixed encoder representations without enabling learning of new visual features. The visual probes proposed in VisCOP differ fundamentally, as they are designed to *extract and learn* domain-specific visual representations rather than to simply *compress* or *condition on* pretrained ones. This is enabled by their interaction with intermediate representations of the vision encoder, allowing the probes to extract domain-specific representations that are not propagated to the final representation of the pretrained vision encoder [37, 53].

3 Problem Formulation

Let $\mathcal{S}$ denote the *source domain*, on which the vision-language model f_{θ^0} has been pretrained, and let $\mathcal{T}$ denote the *target domain*, the domain of interest for adaptation. The two domains differ in their underlying distributions (e.g., viewpoint, modality, or task), which causes f_{θ^0} to perform poorly when directly applied to $\mathcal{T}$.

Training supervision in these domains is provided as video-QA pairs (v, q, a) , where v is a video, q is an instruction or question, and a is the corresponding response. While f_{θ^0} has been pretrained on samples $(v, q, a) \sim \mathcal{S}$, at adaptation time we only assume availability of target domain samples $(v, q, a) \sim \mathcal{T}$. The objective of domain adaptation is to update the pretrained parameters θ^0 to obtain θ^* that improves performance on domain $\mathcal{T}$, while retaining performance on domain $\mathcal{S}$. Formally,

$$R_{\mathcal{T}}(\theta^*) < R_{\mathcal{T}}(\theta^0) \quad \text{and} \quad R_{\mathcal{S}}(\theta^*) \approx R_{\mathcal{S}}(\theta^0)$$

where $R_{\mathcal{D}}$ denotes the VLM’s expected autoregressive next-token prediction loss under domain $\mathcal{D}$. In summary, our problem statement considers adaptation of a pretrained VLM to a novel domain using only video-QA pairs from that domain. The objective is to improve target-domain performance while minimizing catastrophic forgetting of source-domain capabilities. In the next section, we introduce our proposed method, which enables balanced domain adaptation under these constraints.

4 Method: Video Domain-adaptive VLM

Given a video input $\mathbf{V} = \{\mathbf{I}_t\}_{t=1}^T$ consisting of T frames, the goal of the VLM is to generate the response corresponding to the input instruction in an autoregressive manner.

4.1 Preliminary

Existing VLMs for video representation learning [3, 54] consist of three standard components: (i) a vision encoder that maps visual inputs into a sequence of spatio-temporal tokens, (ii) a vision-language connector that projects the visual tokens to the embedding space of a language model, and (iii) an LLM that processes the projected visual tokens jointly with language tokens to enable multi-modal reasoning. For the input video $\mathbf{V}$, each frame $\mathbf{I}_t$ is processed independently by the vision encoder through a stack of L transformer layers. The visual tokens after the ℓ -th layer are denoted as

$$\mathbf{X}_t^\ell \in \mathbb{R}^{N \times d_v}, \quad \ell = 1, \dots, L$$

where N is the number of spatial patch tokens per frame and d_v is the embedding dimension of the vision encoder. Concatenating these tokens over time yields $\mathbf{X}^\ell \in \mathbb{R}^{(TN) \times d_v}$ which represents the sequence of spatio-temporal visual tokens at the ℓ -th layer of the vision encoder. The final layer outputs $\mathbf{X}^L$ are then projected to the language embedding space via a vision-language connector $\mathcal{C}$ to obtain the visual embeddings used as input to the LLM

$$\mathbf{E} = \mathcal{C}(\mathbf{X}^L) \in \mathbb{R}^{(T\tilde{N}) \times d_{\text{lm}}}$$

where $\tilde{N}$ is the number of visual tokens input to the LLM after spatial down-sampling [54], and d_{lm} is the embedding dimension of the LLM.

The VLM is then trained to optimize a standard autoregressive next token prediction loss. Specifically, given the visual embeddings $\mathbf{E}$ and the tokenized QA pair $(\mathbf{Q}, \mathbf{A})$, we optimize the likelihood of predicting $\mathbf{A}$ conditioned on the visual embeddings and the question

$$P(\mathbf{A} \mid \mathbf{E}, \mathbf{Q}) = \prod_{j=1}^{\text{Len}} P_{\theta}(\mathbf{a}_j \mid \mathbf{E}, \mathbf{Q}, \mathbf{A}_{<j})$$

where θ are the trainable parameters of the VLM, Len indicates the token length of $\mathbf{A}$, and $\mathbf{A}_{<j}$ represents the subsequence of answer tokens preceding position j .

For domain-adaptive post training of VLMs, finetuning the vision encoder of a pretrained VLM for a target domain $\mathcal{T}$ often leads to overfitting on $\mathcal{T}$ and catastrophic forgetting of the source domain [25, 49, 51]. To mitigate this trade-off, a domain-adaptive pathway is required that adapts the VLM to $\mathcal{T}$ while retaining performance on $\mathcal{S}$.

4.2 VisCoP: Vision Contextualized Probing

To capture the relevant visual context that would otherwise be lost by freezing the vision encoder, we propose Vision Contextualized Probing (**VisCoP**), a

After the final layer, the updated probes $\mathbf{P}^L$ are projected to the language embedding space via a dedicated connector $\mathcal{C}_{\text{probe}}$,

$$\mathbf{Z} = \mathcal{C}_{\text{probe}}(\mathbf{P}^L) \in \mathbb{R}^{M \times d_{\text{lm}}},$$

and the VLM is trained with the standard autoregressive objective additionally conditioned on $\mathbf{Z}$:

$$P(\mathbf{A} \mid \mathbf{E}, \mathbf{Q}, \mathbf{Z}) = \prod_{j=1}^{\text{Len}} P_{\theta}(\mathbf{a}_j \mid \mathbf{E}, \mathbf{Q}, \mathbf{Z}, \mathbf{A}_{<j}).$$

Thus, the probes act as low-dimensional control knobs that bias learning toward domain-relevant structure and away from spurious artifacts. This is reinforced by applying updates through the probe connector, and through LoRA [17] updates in the LLM embedding space, which confine parameter changes to a low-rank, probe-defined visual subspace that preserves generalizable behavior while enabling targeted specialization.

5 Experiments

We evaluate VisCoP for effective domain adaptation and minimal forgetting. Section 5.1 details the setup (architecture, training, metrics); Section 5.2 reports results on egocentric, depth, and robotic-control targets; Section 5.3 presents ablations and representation analyses of the probes and interaction modules.

5.1 Experimental Setting

VLM Architecture. We consider a VLM architecture consisting of a SigLIP [53] vision encoder, Qwen 2.5 [36] LLM, and a 2-layer MLP vision-language connector, with all modules initialized from the pretrained weights of VideoL-LaMA3 [54]. The embedding dimension of the vision encoder is $d_v = 1152$, and the embedding dimension of the LLM is $d_{\text{lm}} = 3584$. We refer to this pretrained model as the *base* VLM, and to models adapted to a target domain as *expert* VLMs. To adapt the base VLM to a target domain, we perform fine-tuning on the target domain with a learning rate of 1×10^{-5} for the LLM and vision-language connector, and a learning rate of 2×10^{-6} for the vision encoder (when trainable). The model is finetuned on 4 NVIDIA H200 GPUs for 3 epochs when adapting to video domains, or 2 epochs when adapting to robotic control domains.

VisCoP Details. By default, VisCoP operates at every layer of the vision encoder and employs $M = 16$ visual probes unless otherwise stated. The visual probes are initialized from the normal distribution $\mathcal{N}(0, 0.02)$. Each interaction module Φ^ℓ is implemented as a multi-head cross-attention [45], and its weights are initialized from the self-attention weights of the vision encoder at layer ℓ . During domain adaptation, we freeze the vision encoder and update only the visual probes, interaction modules, vision-language connectors, and the LLM’s LoRA parameters. For adaptation to video understanding domains, we update the LLM using LoRA ($r = 16$), while the entire LLM is updated when adapting to the robotic control domain. We keep the vision-language connector trainable throughout, as is standard in VLM training [28, 47].

Adaptation Metrics. We evaluate the domain adaptation of VLMs across two dimensions: (i) their “*improvement*” on the target domain $\mathcal{T}$, and (ii) their “*retention*” on the source domain $\mathcal{S}$. Improvement on the target domain is measured as the performance difference between the expert and base VLMs on target domain benchmarks; retention is the corresponding difference on source domain benchmarks. If $\text{Acc}_{\mathcal{D}}$ denotes the average accuracy over all benchmarks within the domain $\mathcal{D}$, then the metrics are computed by:

$$\Delta_{\text{target}} = \text{Acc}_{\text{target}}^{\text{expert}} - \text{Acc}_{\text{target}}^{\text{base}} \quad \Delta_{\text{source}} = \text{Acc}_{\text{source}}^{\text{expert}} - \text{Acc}_{\text{source}}^{\text{base}}$$

5.2 Source and Target Domains

The source domain $\mathcal{S}$ is fixed throughout this paper: exocentric RGB videos of human actions reflecting the samples used to train generic VLMs for video representation learning. Our target domains $\mathcal{T}$ deliberately shift the input distribution, and consist of (1) egocentric video understanding, (2) depth-modality video understanding, and (3) robotic control. All data (videos and instructions) in our chosen target domain benchmarks were not used in the pretraining of the base VLM [54]. We evaluate ViSCoP’s adaptation to each target while measuring retention of source domain competencies: (i) when adapting to egocentric video, exocentric understanding should be preserved; (ii) when adapting to depth video, RGB understanding should be preserved; and (iii) when adapting to robotic control, human-action understanding should be preserved.

Training datasets. For ego and depth video understanding domains, we adapt using EgoExo4D [14], a large-scale multi-view dataset containing time-synchronized egocentric and exocentric videos of skilled human activities. We utilize a total of 24,688 videos from the keystep recognition subset to generate 74,064 video instruction pairs. These instructions are recaptioned from the instruction pairs provided in [42]. For the **egocentric** target domain, we adapt on 45,888 egocentric video-instruction pairs. For the **depth** target domain, we convert all exocentric RGB videos to depth using DepthAnythingV2 [48] while keeping the language instructions unchanged, yielding 28,176 depth instruction pairs.

We perform adaptation to the **robotic control** domain in both simulated and real-world robot environments. In the *simulated environment*, we leverage the training set of VIMA-Bench [20]. VIMA-Bench contains 17 object manipulation tasks with an action space comprising two 2D coordinates (for pick and place positions) and two quaternions (for rotation). Since the training set of VIMA-Bench lacks natural language instructions by default, we leverage the instruction pairs generated in LLaRA [26], resulting in 13,922 instruction pairs across 7,995 action trajectories. In the *real-world environment*, we collect a dataset using a 6-DoF xArm 7 robot arm deployed in a tabletop manipulation setting. This dataset, which we refer to as xArm-Det, contains 1,007 instruction pairs depicting novel objects and spatial configurations not present in simulation. During adaptation, we train jointly on VIMA-Bench and xArm-Det, resulting in a total of 14,929 instruction pairs. The large-scale simulated data enables the model to learn manipulation skills, while xArm-Det exposes the model to our novel robot

environment. Illustrations of our real-world robot environment and examples from VIMA-Bench are provided in the Appendix.

Table 1: Egocentric Video Understanding Experts. Performance of adaptation strategies on the egocentric target domain and exocentric source domain. 🔥 denotes trainable components, ❄️ denotes frozen components, and 🔥❄️ denotes LoRA updates. Δ_{target} and Δ_{source} denote relative gains over the Base VLM.

Adaptation Strategy		Egocentric (Target)						Exocentric (Source)					Metrics	
VE	LLM	Action Und.	Task Regions	HOI	Hand Ident.	EgoSchema	Avg	NeXTQA	Video MME	ADL MCQ	ADL Desc	Avg	Δ_{target}	Δ_{source}
Base VLM		75.37	74.88	75.56	65.38	60.98	70.43	84.32	65.37	77.36	70.65	74.42	-	-
Finetuning Strategies (🔥 Vision Language Connector)														
❄️	❄️	73.00	76.71	72.85	65.51	60.43	69.70	84.21	62.67	76.56	75.51	74.74	-0.74	+0.31
🔥	❄️	<u>76.13</u>	82.93	73.32	64.86	61.14	<u>71.68</u>	83.87	61.41	77.05	<u>76.09</u>	74.61	+1.24	+0.18
🔥	🔥	73.28	82.68	72.96	65.77	60.31	71.00	82.34	64.26	<u>78.21</u>	70.89	73.93	+0.57	-0.50
❄️	🔥	73.49	74.27	74.50	<u>64.99</u>	<u>61.52</u>	69.75	84.24	<u>64.41</u>	77.42	74.36	<u>75.11</u>	-0.68	+0.68
VisCoP		81.28	82.80	78.75	64.86	62.11	73.96	84.31	64.70	78.97	76.78	76.19	+3.53	+1.77

Egocentric Video Understanding Target and source benchmarks. For evaluation on the **target domain**, we evaluate on the Ego-in-Exo Perception [42] and EgoSchema [32] benchmarks. Ego-in-Exo Perception is derived from EgoExo4D and comprises 3,991 video question-answer (video-QA) pairs spanning four categories: action understanding (Action Und.), task-relevant region understanding (Task Regions), human-object interactions (HOI), and hand identification (Hand Ident.). Because it is derived from EgoExo4D, Ego-in-Exo Perception can be evaluated from either the egocentric or the exocentric viewpoint. For the ego target domain experiments, we report results using the egocentric videos, denoted as Ego-in-Exo Perception (Ego RGB). EgoSchema consists of 5,031 egocentric video-QA pairs derived from Ego4D [13].

For evaluation on the **source domain**, we select benchmarks that measure exocentric video understanding capability. Specifically, we evaluate on the NeXTQA [46], VideoMME [11], and ADL-X [41] benchmarks. NeXTQA and VideoMME are general-purpose video-QA benchmarks built from web-scraped videos (e.g., from YouTube), with 8,564 QA pairs in NeXTQA and 2,700 QA pairs in VideoMME. ADL-X is a video-QA benchmark built from videos of activities of daily living, it contains a total of 10,561 multiple-choice questions (ADL-X MCQ) and 1,862 video description questions (ADL-X Desc) derived from various activities of daily living datasets [8, 9, 18, 44].

Results. Table 1 reports results of adaptation to the egocentric viewpoint. Training only the vision-language connector or the connector together with LLM LoRA adapters does not lead to effective adaptation to the target domain ($\Delta_{\text{target}} < 1$). Updating all three modules (connector, vision encoder, and LLM) improves performance on the target domain by $\Delta_{\text{target}} = +0.57$, but the large number of trainable parameters results in forgetting on the source benchmarks ($\Delta_{\text{source}} = -0.50$). In contrast, updating the connector and vision encoder alone slightly improves performance on the target domain and does not lead to forgetting on the source domain. These results highlight that the core difficulty of domain adaptation in existing VLMs arises from the necessity of updating the vision encoder to learn domain-specific visual representations, which inevitably

leads to forgetting of pretrained knowledge. Our proposed VisCoP achieves the strongest adaptation performance, with the highest improvement on the target domain ($\Delta_{\text{target}} = +3.5$) while simultaneously maintaining retention on the source benchmarks ($\Delta_{\text{source}} = +1.8$). Interestingly, VisCoP not only avoids catastrophic forgetting but also improves performance on some source benchmarks (e.g., ADL-X). We attribute this positive transfer to a multi-axis domain shift: although source and target differ in viewpoint (exocentric vs. egocentric), their action distributions overlap. ADL-X, while exocentric, encapsulates activities of daily living that closely aligns with the EgoExo4D action distribution, enabling beneficial cross-domain generalization.

Table 2: Depth Video Understanding Experts. Performance of adaptation strategies on the depth target domain and RGB source domain. 🔥 denotes trainable components, ❄️ denotes frozen components, and 🔥❄️ denotes LoRA updates. Δ_{target} and Δ_{source} denote relative gains over the Base VLM.

Adaptation Strategy		Depth (Target)					RGB (Source)						Metrics	
VE	LLM	Action Und.	Task Regions	HOI	Hand Ident.	Avg	Ego-in-Exo (Exo RGB)	NeXTQA	Video MME	ADL MCQ	ADL Desc	Avg	Δ_{target}	Δ_{source}
Base VLM		34.73	50.61	35.06	63.06	45.86	66.27	84.32	65.37	77.36	70.65	<u>72.79</u>	–	–
Finetuning Strategies (🔥 Vision Language Connector)														
❄️	❄️	55.67	66.59	<u>62.46</u>	64.49	<u>62.30</u>	<u>71.36</u>	83.15	62.41	70.90	69.05	71.37	+16.44	-1.42
🔥	❄️	57.20	<u>69.63</u>	54.43	64.48	61.44	60.97	82.89	62.00	71.48	67.26	68.92	+15.57	-3.87
❄️	🔥❄️	42.94	53.54	43.92	63.96	51.09	60.97	83.73	64.19	72.19	<u>72.49</u>	70.71	+5.23	-2.08
VisCoP		<u>56.78</u>	73.17	66.23	<u>64.35</u>	65.13	71.89	<u>83.91</u>	<u>64.30</u>	<u>76.59</u>	76.47	74.63	+19.27	+1.84

Depth Video Understanding Target and source benchmarks. For evaluation on the **target domain**, we evaluate on the Ego-in-Exo Perception [42] benchmark. In the depth-adaptation setting, we train on depth maps of exocentric videos extracted with DepthAnythingV2 [48] and evaluate on exocentric depth videos following [42], denoted Ego-in-Exo Perception (Exo Depth). For the **source domain**, we use *RGB* benchmarks of exocentric understanding: Ego-in-Exo Perception (Exo RGB), NeXTQA, VideoMME, and ADL-X.

Results. We present the results for adaptation to the depth modality in Table 2. In contrast to the results on egocentric viewpoint adaptation, we find that all training strategies achieve improvements on the target domain, reflecting the disparity of the visual embedding space between the depth and RGB modalities. We find that this disparity leads to different behavior across training strategies. Jointly updating the vision encoder and the vision-language connector preserves source performance for egocentric adaptation but causes severe catastrophic forgetting under depth adaptation ($\Delta_{\text{source}} = -3.87$). This arises from the substantial encoder updates required to bridge RGB and depth, which overwrite RGB representations. In contrast, VisCoP preserves RGB features and source performance while achieving the largest target domain gains ($\Delta_{\text{target}} = +19.27$).

Robot Control Target and source benchmarks. For evaluation on the **target domain**, we consider both simulated and real-world robotic environments. In simulation, we use the evaluation set of VIMA-Bench [20], which organizes tasks into three levels of difficulty: L1 (Object Placement), where all

Table 3: Robot Control Experts (Simulation). Performance of adaptation strategies on the robotic control target domain and human understanding source domain. 🔥 denotes trainable components, ❄ denotes frozen components. Δ_{target} and Δ_{source} denote relative gains over the Base VLM.

Adaptation Strategy		Robotic Control (Target)				Human Understanding (Source)						Metrics	
VE	LLM	L1	L2	L3	Avg	Ego-in-Exo (Exo RGB)	NeXTQA	Video MME	ADL MCQ	ADL Desc	Avg	Δ_{target}	Δ_{source}
Base VLM		0	0	0	0	66.27	84.32	65.37	77.36	70.65	72.79	—	—
Finetuning Strategies (🔥 Vision Language Connector)													
🔥	🔥	69.62	60.77	65.00	65.13	56.92	83.24	62.74	52.21	64.50	63.92	+65.13	-8.87
🔥	❄	63.46	63.08	68.75	65.10	59.42	83.16	64.41	52.92	64.86	64.95	+65.10	-7.84
VisCoP		67.69	65.77	70.00	67.82	71.19	83.71	63.67	55.89	66.62	68.22	+67.82	-4.58

objects have been seen during training; L2 (Novel Combination), where objects seen during training appear in new pairings or contexts; and L3 (Novel Objects), where objects entirely unseen during training are introduced. Together, these levels measure generalization from familiar training conditions to progressively more challenging distributions. In the real-world setting, we evaluate on three tabletop manipulation tasks: T1) Place the {object} on the plate, T2) Pick up and rotate {object} by {angle}; and T3) Move all {color} objects onto the plate. Examples of each task and a list of objects used is provided in the Appendix. On these robotic control benchmarks, the reported accuracy corresponds to the success rate across all robot manipulation tasks. For **source domain** evaluation of VLMs trained on both real and simulated robotic environments, we use the human-activity video benchmarks Ego-in-Exo (Exo RGB), NeXTQA, VideoMME, and ADL-X.

Results. The results of adaptation to the robotic control domain are presented in Table 3. The base VLM demonstrates weak performance on all robot control tasks, as its pretraining distribution lacks action trajectories (i.e., instruction data mapping from visual observations to robot actions). This lack of pretraining results in 0% accuracy across all levels of VIMA-Bench, and is consistent with prior works [26]. This highlights the extreme domain gap both in the visual space (robot observations vs. human videos) and in the language space (control actions vs. linguistic outputs) between the source and target domains. Similarly to the depth adaptation setting, we find that training the vision encoder improves performance on the target domain, but results in the worst source domain retention ($\Delta_{\text{source}} = -8.87$) of all robot control experts. In contrast, our proposed VisCoP achieves the best performance on the target domain ($\Delta_{\text{target}} = +67.82$) while retaining the most source domain knowledge ($\Delta_{\text{source}} = -4.58$) compared to other experts, demonstrating the effectiveness of our method even when the gap between the source and target domains is very large. Also note that Vis-

Table 4: Robot Control Experts (Real-world). Performance on the robotic control target domain and human understanding source domain. 🔥 denotes trainable components and ❄ denotes frozen components.

Adaptation Strategy		Robotic Control (Target)				Metrics	
VE	LLM	T1	T2	T3	Avg	Δ_{target}	Δ_{source}
Training data: VIMA-Bench							
🔥	🔥	45.00	60.00	15.00	40.00	+40.00	-8.87
VisCoP		40.00	70.00	20.00	43.33	+43.33	-4.58
Training data: VIMA-Bench + xArm-Det							
🔥	🔥	85.00	85.00	70.00	80.00	+80.00	-11.04
VisCoP		100.00	100.00	90.00	96.67	+96.67	-11.00

Table 5: Comparison with other adaptation strategies. Legend: *VP* (visual probing with no interaction modules), *Last-4* (train only the last 4 vision encoder layers), *QFormer Style* (interaction modules placed after the last VE layer).

Method	Target	Source	Adaptation Metrics	
	Avg	Avg	Δ_{target}	Δ_{source}
Base VLM	70.43	74.42	—	—
<i>VP</i>	65.57	75.05	-4.86	+0.62
LoRA (Full)	69.85	75.35	-0.59	+0.92
<i>Last-4</i>	70.46	72.62	+0.02	-1.80
<i>VPT</i> [19]	71.36	74.97	+0.93	+0.55
<i>QFormer Style</i>	70.99	75.03	+0.56	+0.61
Model Tailor [58]	70.27	75.29	-0.16	+0.86
VisCoP	73.96	76.19	+3.53	+1.77

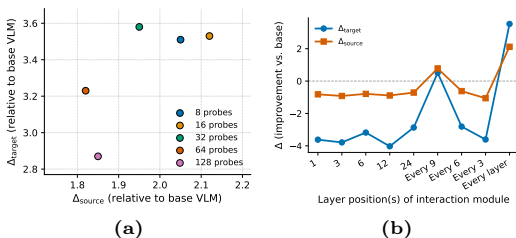


Fig. 3: Ablation studies of VisCoP on the egocentric target domain. (a) Effects of varying the number of visual probes within VisCoP (default=16). (b) Effect of interaction module placement position (default=Every layer).

CoP operates on per-timestep images in these experiments; thus the visual probes consume the same visual tokens as the vision encoder, suggesting they extract domain-specific representations more effectively than the base encoder.

We further evaluate adaptation in the real-world setting using the xArm-Det dataset in Table 4. We consider a *transfer setting*, where the experts are trained only on VIMA-Bench and directly evaluated on xArm-Det, and the setting where the experts are jointly trained on both VIMA-Bench and xArm-Det. In both cases, our proposed VisCoP outperforms the vision encoder trained experts on target domain adaptation as well as source domain retention.

5.3 Model Diagnosis and Analysis

In this section, we motivate the design of VisCoP through a diagnostic study, and perform an analysis on the visual representations it learns. We investigate the number of visual probes, as well as the placement of interaction modules within the vision encoder. We then analyze the domain-specific representations learned by VisCoP through t-SNE and attention visualizations.

Alternatives to learnable queries. Table 5 compares VisCoP against alternative adaptation strategies. *Visual Probes Only (VP)* trains only visual probes with their vision-language connector ($\mathcal{C}_{\text{probe}}$) without any interaction modules. *Partial Encoder Training (Last-4)* makes the final four layers of the vision encoder trainable. *QFormer-Style Compression* uses visual probes with interaction modules only at the vision encoder’s final layer, mimicking Q-Former’s compression approach [24]. Notably, several of these alternatives underperform even the frozen base VLM on target tasks. We attribute this to their reliance on the pre-trained visual pathway: they can recombine existing representations but struggle to introduce novel cues (e.g., egocentric or depth structure) that are absent from those representations. Full fine-tuning is more competitive on the target precisely because it can overwrite this pathway, at the cost of source-domain performance. VisCoP occupies the middle ground, introducing a parallel probing pathway that decouples target-domain learning from the frozen pretrained

pathway, and thereby learns new cues without overwriting the source representation. *Model Tailor* [58] performs post-hoc domain adaptation by fusing parameter updates from a fine-tuned VLM back into the base VLM, modifying only the LLM parameters and leaving the vision encoder untouched. Training with QFormer-Style compression or only training with visual probes (VP) underperforms compared to VisCoP, indicating the importance of probe interactions at intermediate layers of the vision encoder to learn domain-specific features across multiple levels of abstraction. Similarly, training only the last four layers of the vision encoder, or training it with LoRA, also underperforms, highlighting that even partial parameter updates fail to capture domain-specific signals as effectively as VisCoP’s visual probes. Model Tailor also falls short in this setting, suggesting that approaches which do not leverage intermediate vision encoder representations struggle to learn domain-specific visual features.

Table 6: Ablation on probe interaction scope. Spatial-only probing restricts probe interactions within individual frames, while temporal-only probing restricts interactions across time. Full VisCoP jointly models spatial and temporal structure.

Interaction Module Scope	Target	Source	Δ_{target} (↑)	Δ_{source} (↑)
Base VLM	70.43	74.42	–	–
Spatial Only	72.26	76.12	+1.83	+1.70
Temporal Only	70.13	75.52	-0.30	+1.10
Spatio-Temporal (VisCoP)	73.96	76.19	+3.53	+1.77

indicates that domain-relevant signals emerge from the joint spatio-temporal structure of video representations, and that enabling probes to integrate information across both dimensions is critical for effective video domain adaptation.

Computation overhead of VisCoP.

VisCoP introduces only modest computational overhead relative to the base VLM, introducing only 2% more parameters than the base VLM. In Table 7, we compute the average VRAM usage and latency

during inference on the Ego-in-Exo Perception benchmark, as well as the total parameter count. We find that VisCoP increases VRAM usage by 3.4GB and adds just 0.013s of inference latency in visual feature extraction (Before LLM). Interestingly, the inference latency of our model is lower than that of the base VLM. We attribute this to the base VLM producing longer, less focused responses, which increases total decoding time.

Ablations on probes and interaction modules. We study the effect of the number of visual probes and the placement of interaction modules (Figure 2a, Figure 2b). Probes consistently improve performance over the base VLM, with the best trade-off at 16 probes ($\Delta_{\text{target}} = +3.53$, $\Delta_{\text{source}} = +2.12$); larger probe counts offer no further gains and can reduce performance due to redundancy. For

Effect of interaction module scope.

A key design choice in VisCoP is the scope of visual features that the probes interact through the interaction modules. In Table 6, restricting probes to spatial-only or temporal-only interactions yields weaker adaptation performance, while full spatio-temporal interaction achieves substantially larger gains ($\Delta_{\text{target}} = +3.53$). This result in-

Table 7: Computation overhead of VisCoP.

Model	Max VRAM	Latency (Full Model)	Latency (Before LLM)	Total Num Params.
Base VLM	24.4GB	0.767s	0.056s	8.04B
VisCoP	27.8GB	0.417s	0.069s	8.21B

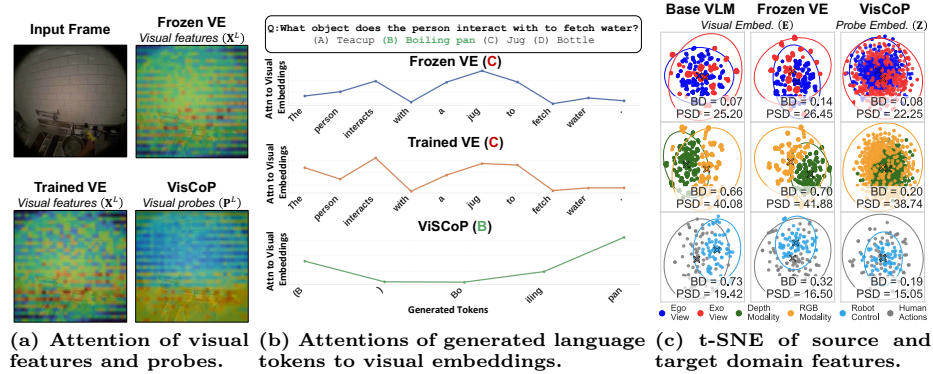


Fig. 4: Analysis of VisCoP. (a) Attentions between visual features and visual probes. (b) Attention of generated language tokens to visual embeddings. (c) t-SNE of visual and probe embeddings. Ellipses denote 95% confidence regions of a fitted 2D Gaussian, and cross markers indicate the Gaussian means. Bhattacharyya distance (BD) and per-sample distance (PSD) are shown.

interaction modules, applying them at every encoder layer yields the strongest adaptation, while sparse placement (e.g., every 6 or 9 layers) provides weaker or inconsistent gains. These results highlight the importance of using a small number of probes with dense access to intermediate features.

Visualizing attention in domain-adapted VLMs In Figure 4a, we analyze attention maps of various VLM adaptation strategies to assess how different components capture domain-specific visual features. For both the frozen and trainable vision encoders, we visualize attention using attention rollout [1], for VisCoP we visualize the attentions of the visual probes, averaged across all probes. The frozen vision encoder fails to focus consistently on relevant regions under the experimented domains, reflecting its limited ability to capture domain-specific features. The trained vision encoder yields sharper attention on the relevant regions, indicating its ability to learn domain-specific features, albeit at the cost of catastrophic forgetting of the source domain as shown in Section 5.2. In contrast, the visual probes of VisCoP have a sharp focus on the task-relevant regions, despite the vision encoder being frozen. This indicates that the probes alone are able to extract the domain-specific visual features necessary for adaptation. In Figure 4b, we visualize the attention of generated language tokens to visual embeddings. We find that VisCoP correctly responds to the query, with more focus given to tokens corresponding to relevant objects.

Learning domain-specific representations. Figure 4c compares t-SNE embeddings of source and target domains across different VLMs. Circles represent individual samples, and ellipses denote 95% confidence regions of fitted 2D Gaussians. For the egocentric and depth target domains, each source-target pair corresponds to time-synchronized videos of the same action. For the robot target domain, pairs correspond to pick-and-place actions performed by humans. Ideally, the embeddings of paired samples should lie closer together in the embedding space, reflecting alignment across the source and target domains. We quantify

this using two metrics: the *Bhattacharyya distance* (BD) computed between the Gaussians fitted to each domain, and the *per-sample distance* (PSD), defined as the mean Euclidean distance between paired embeddings across domains. We observe that the visual probes of VisCoP learn stronger alignment between the source and target domains. Quantitatively, VisCoP yields the smallest Bhattacharyya distance and per-sample distance between paired source and target embeddings across all three target domains, indicating tighter cross-domain alignment than both the base VLM and the frozen-encoder baseline. The effect is most pronounced under the largest shifts (depth and robot control), where the frozen encoder’s paired embeddings remain widely separated while the probes draw them into closer correspondence.

Table 8: Audio Understanding Experiments. Adaptation to a non-visual (spectrogram) target domain on MusicAVQA, with video understanding as the source domain.

Method	Target	Avg Source	Avg Δ_{target}	Δ_{source}
Base VLM	28.3	57.5	-	-
UE + LLM (LoRA)	39.8	52.2	+11.5	-5.3
VisCoP	40.2	56.8	+11.9	-0.7

Generalization beyond visual domains. The target domains considered thus far all shift the visual input distribution. To test whether VisCoP generalizes to non-visual shifts, we adapt to *audio understanding*, where audio is represented as a spectrogram, using video understanding as

the source domain. Following OneLLM [16], we repurpose the vision encoder into a universal encoder (UE) that processes both audio and video, and instantiate VisCoP on top of it. We adapt and evaluate on MusicAVQA [23], treating its audio-centric questions as the target domain and its video-centric questions as the source. As shown in Table 8, VisCoP and the UE + LLM (LoRA) baseline are comparable on the target domain ($\Delta_{\text{target}} = +11.9$ vs. $+11.5$), but VisCoP nearly eliminates source-domain forgetting ($\Delta_{\text{source}} = -0.7$ vs. -5.3). This demonstrates VisCoP is not specific to visual shifts: it extracts domain-relevant cues while preventing forgetting even when the target modality is non-visual.

6 Conclusion

We introduced VisCoP, a mechanism that extracts domain-specific visual features through probing of a frozen vision encoder to enable effective domain adaptation in VLMs and prevent catastrophic forgetting. VLMs equipped with VisCoP achieve superior target domain performance, while maintaining strong source domain capabilities across cross-view, cross-modal, and cross-task adaptation scenarios. We will release all code, models, and evaluation protocols to facilitate future research.

7 Acknowledgments

This work was supported in part by the National Science Foundation (IIS-2245652) and the University of North Carolina at Charlotte. Computational resources were provided by the NSF National AI Research Resource Pilot (NAIRR-240338) and the NCShare initiative. We would like to thank Hieu Le and other members of the Charlotte Vision Lab for their valuable insights and discussions.

References

1. Abnar, S., Zuidema, W.: Quantifying attention flow in transformers. In: Annual Meeting of the Association for Computational Linguistics (2020) **14**
2. Alayrac, J.B., Donahue, J., Luc, P., Miech, A., Barr, I., Hasson, Y., Lenc, K., Mensch, A., Millican, K., Reynolds, M., Ring, R., Rutherford, E., Cabi, S., Han, T., Gong, Z., Samangooei, S., Monteiro, M., Menick, J.L., Borgeaud, S., Brock, A., Nematzadeh, A., Sharifzadeh, S., Bińkowski, M., Barreira, R., Vinyals, O., Zisserman, A., Simonyan, K.: Flamingo: a visual language model for few-shot learning. In: Koyejo, S., Mohamed, S., Agarwal, A., Belgrave, D., Cho, K., Oh, A. (eds.) Advances in Neural Information Processing Systems. vol. 35, pp. 23716–23736. Curran Associates, Inc. (2022), https://proceedings.neurips.cc/paper_files/paper/2022/file/960a172bc7fbf0177ccccbb411a7d800-Paper-Conference.pdf **2, 4, 6**
3. Bai, S., Chen, K., Liu, X., Wang, J., Ge, W., Song, S., Dang, K., Wang, P., Wang, S., Tang, J., Zhong, H., Zhu, Y., Yang, M., Li, Z., Wan, J., Wang, P., Ding, W., Fu, Z., Xu, Y., Ye, J., Zhang, X., Xie, T., Cheng, Z., Zhang, H., Yang, Z., Xu, H., Lin, J.: Qwen2.5-vl technical report. arXiv preprint arXiv:2502.13923 (2025) **1, 5**
4. Bertasius, G., Wang, H., Torresani, L.: Is space-time attention all you need for video understanding? In: Proceedings of the International Conference on Machine Learning (ICML) (July 2021) **2**
5. Chen, J., Gui, C., Ouyang, R., Gao, A., Chen, S., Chen, G.H., Wang, X., Zhang, R., Cai, Z., Ji, K., Yu, G., Wan, X., Wang, B.: Huatuoogpt-vision, towards injecting medical visual knowledge into multimodal llms at scale (2024) **3**
6. Chen, L., Wei, X., Li, J., Dong, X., Zhang, P., Zang, Y., Chen, Z., Duan, H., Lin, B., Tang, Z., Yuan, L., Qiao, Y., Lin, D., Zhao, F., Wang, J.: Sharegpt4video: Improving video understanding and generation with better captions. In: Advances in Neural Information Processing Systems (2024) **1**
7. Cheng, D., Huang, S., Zhu, Z., Zhang, X., Zhao, W.X., Luan, Z., Dai, B., Zhang, Z.: On domain-adaptive post-training for multimodal large language models. In: Conference on Empirical Methods in Natural Language Processing Findings (2025) **3**
8. Dai, R., Das, S., Sharma, S., Minciullo, L., Garattoni, L., Bremond, F., Francesca, G.: Toyota smarthome untrimmed: Real-world untrimmed videos for activity detection. IEEE Transactions on Pattern Analysis and Machine Intelligence (2022). <https://doi.org/10.1109/TPAMI.2022.3169976> **9**
9. Das, S., Dai, R., Koperski, M., Minciullo, L., Garattoni, L., Bremond, F., Francesca, G.: Toyota smarthome: Real-world activities of daily living. In: Proceedings of the IEEE/CVF International Conference on Computer Vision. pp. 833–842 (2019) **9**
10. Evci, U., Dumoulin, V., Larochelle, H., Mozer, M.C.: Head2toe: Utilizing intermediate representations for better transfer learning. In: Proceedings of the 39th International Conference on Machine Learning (ICML) (2022) **3**
11. Fu, C., Dai, Y., Luo, Y., Li, L., Ren, S., Zhang, R., Wang, Z., Zhou, C., Shen, Y., Zhang, M., Chen, P., Li, Y., Lin, S., Zhao, S., Li, K., Xu, T., Zheng, X., Chen, E., Shan, C., He, R., Sun, X.: Video-mme: The first-ever comprehensive evaluation benchmark of multi-modal large language models in video analysis. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (2025) **9**

12. Gao, P., Geng, S., Zhang, R., Ma, T., Fang, R., Zhang, Y., Li, H., Qiao, Y.: Clip-adapter: Better vision-language models with feature adapters. In: *International Journal of Computer Vision* (2023) 3
13. Grauman, K., Westbury, A., Byrne, E., Chavis, Z., Furnari, A., Girdhar, R., Hamburger, J., Jiang, H., Liu, M., Liu, X., Martin, M., Nagarajan, T., Radosavovic, I., Ramakrishnan, S.K., Ryan, F., Sharma, J., Wray, M., Xu, M., Xu, E.Z., Zhao, C., Bansal, S., Batra, D., Cartillier, V., Crane, S., Do, T., Doulaty, M., Erapalli, A., Feichtenhofer, C., Fragomeni, A., Fu, Q., Gebreselasie, A., Gonzalez, C., Hillis, J., Huang, X., Huang, Y., Jia, W., Khoo, W., Kolar, J., Kottur, S., Kumar, A., Landini, F., Li, C., Li, Y., Li, Z., Mangalam, K., Modhugu, R., Munro, J., Murrell, T., Nishiyasu, T., Price, W., Ruiz Puentes, P., Ramazanova, M., Sari, L., Somasundaram, K., Southerland, A., Sugano, Y., Tao, R., Vo, M., Wang, Y., Wu, X., Yagi, T., Zhao, Z., Zhu, Y., Arbelaez, P., Crandall, D., Damen, D., Farinella, G.M., Fuegen, C., Ghanem, B., Ithapu, V.K., Jawahar, C.V., Joo, H., Kitani, K., Li, H., Newcombe, R., Oliva, A., Park, H.S., Rehg, J.M., Sato, Y., Shi, J., Shou, M.Z., Torralba, A., Torresani, L., Yan, M., Malik, J.: Ego4d: Around the world in 3,000 hours of egocentric video. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 18995–19012 (2022) 9
14. Grauman, K., Westbury, A., Torresani, L., Kitani, K., Malik, J., Afouras, T., Ashutosh, K., Baiyya, V., Bansal, S., Boote, B., Byrne, E., Chavis, Z., Chen, J., Cheng, F., Chu, F.J., Crane, S., Dasgupta, A., Dong, J., Escobar, M., Forigua, C., Gebreselasie, A., Hareesh, S., Huang, J., Islam, M.M., Jain, S., Khirrodar, R., Kukreja, D., Liang, K.J., Liu, J.W., Majumder, S., Mao, Y., Martin, M., Mavroudi, E., Nagarajan, T., Ragusa, F., Ramakrishnan, S.K., Seminara, L., Somayazulu, A., Song, Y., Su, S., Xue, Z., Zhang, E., Zhang, J., Castillo, A., Chen, C., Fu, X., Furuta, R., Gonzalez, C., Gupta, P., Hu, J., Huang, Y., Huang, Y., Khoo, W., Kumar, A., Kuo, R., Lakhavani, S., Liu, M., Luo, M., Luo, Z., Meredith, B., Miller, A., Oguntola, O., Pan, X., Peng, P., Pramanick, S., Ramazanova, M., Ryan, F., Shan, W., Somasundaram, K., Song, C., Southerland, A., Tateno, M., Wang, H., Wang, Y., Yagi, T., Yan, M., Yang, X., Yu, Z., Zha, S.C., Zhao, C., Zhao, Z., Zhu, Z., Zhuo, J., Arbelaez, P., Bertasius, G., Crandall, D., Damen, D., Engel, J., Farinella, G.M., Furnari, A., Ghanem, B., Hoffman, J., Jawahar, C.V., Newcombe, R., Park, H.S., Rehg, J.M., Sato, Y., Savva, M., Shi, J., Shou, M.Z., Wray, M.: Ego-exo4d: Understanding skilled human activity from first- and third-person perspectives. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition* (2025) 8
15. Ha, C.N., Asaadi, S., Karn, S.K., Farri, O., Heimann, T., Runkler, T.: Fusion of domain-adapted vision and language models for medical visual question answering. In: *Proceedings of the Clinical Natural Language Processing Workshop at the 2024 Conference of the North American Chapter of the Association for Computational Linguistics (NAACL)* (2024) 2, 4
16. Han, J., Gong, K., Zhang, Y., Wang, J., Zhang, K., Lin, D., Qiao, Y., Gao, P., Yue, X.: Onellm: One framework to align all modalities with language. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)* (2024) 15
17. Hu, E.J., Shen, Y., Wallis, P., Allen-Zhu, Z., Li, Y., Wang, S., Wang, L., Chen, W.: LoRA: Low-rank adaptation of large language models. In: *International Conference on Learning Representations* (2022), <https://openreview.net/forum?id=nZeVKeeFYf9> 7

18. Jia, B., Chen, Y., Huang, S., Zhu, Y., Zhu, S.C.: Lemma: A multi-view dataset for learning multi-agent multi-task activities. In: Proceedings of the European Conference on Computer Vision. pp. 767–783 (2020) [9](#)
19. Jia, M., Tang, L., Chen, B.C., Cardie, C., Belongie, S., Hariharan, B., Lim, S.N.: Visual prompt tuning. In: ECCV (2022) [2, 4, 6, 12](#)
20. Jiang, Y., Gupta, A., Zhang, Z., Wang, G., Dou, Y., Chen, Y., Fei-Fei, L., Anandkumar, A., Zhu, Y., Fan, L.: Vima: General robot manipulation with multimodal prompts. In: International Conference on Machine Learning (2023) [8, 10](#)
21. Lai, X., Tian, Z., Chen, Y., Li, Y., Yuan, Y., Liu, S., Jia, J.: Lisa: Reasoning segmentation via large language model. arXiv preprint arXiv:2308.00692 (2023) [1](#)
22. Li, C., Wong, C., Zhang, S., Usuyama, N., Liu, H., Yang, J., Naumann, T., Poon, H., Gao, J.: Llava-med: Training a large language-and-vision assistant for biomedicine in one day. In: Advances in Neural Information Processing Systems (2023) [2, 3](#)
23. Li, G., Wei, Y., Tian, Y., Xu, C., Wen, J.R., Hu, D.: Learning to answer questions in dynamic audio-visual scenarios. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) (2022) [15](#)
24. Li, J., Li, D., Savarese, S., Hoi, S.C.H.: Blip-2: Bootstrapping language-image pre-training with frozen image encoders and large language models. In: International Conference on Machine Learning (2023), <https://api.semanticscholar.org/CorpusID:256390509> [2, 4, 6, 12](#)
25. Li, M., Zhong, J., Li, C., Li, L., Lin, N., Sugiyama, M.: Vision-language model fine-tuning via simple parameter-efficient modification. In: Conference on Empirical Methods in Natural Language Processing (2024) [2, 5](#)
26. Li, X., Mata, C., Park, J., Kahatapitiya, K., Jang, Y.S., Shang, J., Ranasinghe, K., Burgert, R., Cai, M., Lee, Y.J., Ryoo, M.S.: Llara: Supercharging robot learning data for vision-language policy. In: International Conference on Learning Representations (2025) [8, 11](#)
27. Li, X.L., Liang, P.: Prefix-tuning: Optimizing continuous prompts for generation. In: ACL (2021) [4](#)
28. Liu, H., Li, C., Wu, Q., Lee, Y.J.: Visual instruction tuning. In: Advances in Neural Information Processing Systems (NeurIPS) (2023) [3, 7](#)
29. Liu, Y., Duan, H., Zhang, Y., Li, B., Zhang, S., Zhao, W., Yuan, Y., Wang, J., He, C., Liu, Z., Chen, K., Lin, D.: Mmbench: Is your multi-modal model an all-around player? In: European Conference on Computer Vision (2024) [1](#)
30. Liu, Z., Ning, J., Cao, Y., Wei, Y., Zhang, Z., Lin, S., Hu, H.: Video swin transformer. 2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) pp. 3192–3201 (2021) [2](#)
31. Maaz, M., Rasheed, H., Khan, S., Khan, F.: Videogpt+: Integrating image and video encoders for enhanced video understanding (2024) [1](#)
32. Mangalam, K., Akshulakov, R., Malik, J.: Egoschema: A diagnostic benchmark for very long-form video language understanding. In: Proceedings of the Thirty-seventh Conference on Neural Information Processing Systems, Datasets and Benchmarks Track (2023) [9](#)
33. Meta: The llama 3 herd of models (2024) [1](#)
34. Mohbat, F., Zaki, M.J.: Llava-chef: A multi-modal generative model for food recipes. In: ACM International Conference on Information and Knowledge Management (2024) [3](#)
35. OpenAI: Thinking with images (April 2025), <https://openai.com/index/thinking-with-images/>, accessed: June 30, 2026 [1](#)

36. Qwen Team, Yang, A., Yang, B., Zhang, B., Hui, B., Zheng, B., Yu, B., Li, C., Liu, D., Huang, F., Wei, H., Lin, H., Yang, J., Tu, J., Zhang, J., Yang, J., Yang, J., Zhou, J., Lin, J., Dang, K., Lu, K., Bao, K., Yang, K., Yu, L., Li, M., Xue, M., Zhang, P., Zhu, Q., Men, R., Lin, R., Li, T., Tang, T., Xia, T., Ren, X., Ren, X., Fan, Y., Su, Y., Zhang, Y., Wan, Y., Liu, Y., Cui, Z., Zhang, Z., Qiu, Z.: Qwen2.5 technical report (2025), <https://arxiv.org/abs/2412.15115> 1, 7
37. Radford, A., Kim, J.W., Hallacy, C., Ramesh, A., Goh, G., Agarwal, S., Sastry, G., Askell, A., Mishkin, P., Clark, J., Krueger, G., Sutskever, I.: Learning transferable visual models from natural language supervision (2021) 1, 3, 4
38. Ranasinghe, K., Shukla, S.N., Poursaeed, O., Ryoo, M.S., Lin, T.Y.: Learning to localize objects improves spatial reasoning in visual-llms. 2024 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) pp. 12977–12987 (2024), <https://api.semanticscholar.org/CorpusID:269043025> 1
39. Rasheed, H., Khattak, M.U., Maaz, M., Khan, S., Khan, F.S.: Finetuned clip models are efficient video learners. In: The IEEE/CVF Conference on Computer Vision and Pattern Recognition (2023) 3
40. Rawal, R., Saifullah, K., Basri, R., Jacobs, D., Somepalli, G., Goldstein, T.: Cinepile: A long video question answering dataset and benchmark (2024) 1
41. Reilly, D., Chakraborty, R., Sinha, A., Govind, M.K., Wang, P., Bremond, F., Xue, L., Das, S.: Llavidal: A large language-vision model for daily activities of living. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (2025) 3, 9
42. Reilly, D., Govind, M.K., Xue, L., Das, S.: From my view to yours: Ego-augmented learning in large vision language models for understanding exocentric daily living activities (2025) 8, 9, 10
43. Ryoo, M.S., Zhou, H., Kendre, S., Qin, C., Xue, L., Shu, M., Park, J., Ranasinghe, K., Savarese, S., Xu, R., Xiong, C., Niebles, J.C.: xgen-mm-vid (blip-3-video): You only need 32 tokens to represent a video even in vlms (2025) 2, 4
44. Sigurdsson, G.A., Varol, G., Wang, X., Farhadi, A., Laptev, I., Gupta, A.: Hollywood in homes: Crowdsourcing data collection for activity understanding. In: Proceedings of the European Conference on Computer Vision. pp. 510–526 (2016) 9
45. Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A.N., Kaiser, L., Polosukhin, I.: Attention is all you need. In: Advances in Neural Information Processing Systems (2017) 2, 7
46. Xiao, J., Shang, X., Yao, A., Chua, T.S.: Next-qa: Next phase of question-answering to explaining temporal actions. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 9777–9786 (2021) 9
47. Xue, L., Shu, M., Awadalla, A., Wang, J., Yan, A., Purushwalkam, S., Zhou, H., Prabhu, V., Dai, Y., Ryoo, M.S., Kendre, S., Zhang, J., Tseng, S., Lujan-Moreno, G.A., Olson, M.L., Hinck, M., Cobbley, D., Lal, V., Qin, C., Zhang, S., Chen, C.C., Yu, N., Tan, J., Awalgaonkar, T.M., Heinecke, S., Wang, H., Choi, Y., Schmidt, L., Chen, Z., Savarese, S., Niebles, J.C., Xiong, C., Xu, R.: xgen-mm (blip-3): A family of open large multimodal models (2025) 1, 7
48. Yang, L., Kang, B., Huang, Z., Zhao, Z., Xu, X., Feng, J., Zhao, H.: Depth anything v2. In: Advances in Neural Information Processing Systems (2024) 8, 10
49. Yang, T., Zhu, Y., Xie, Y., Zhang, A., Chen, C., Li, M.: Aim: Adapting image models for efficient video understanding. In: International Conference on Learning Representations (2023), https://openreview.net/forum?id=CIoSZ_HKHS7 2, 3, 5

50. Yao, H., Zhang, R., Xu, C.: Visual-language prompt tuning with knowledge-guided context optimization. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (2023) [3](#)
51. Zang, Y., Goh, H., Susskind, J., Huang, C.: Overcoming the pitfalls of vision-language model finetuning for ood generalization. In: International Conference on Learning Representations (2024) [2](#), [5](#)
52. Zeng, A., Attarian, M., brian ichter, Choromanski, K.M., Wong, A., Welker, S., Tombari, F., Purohit, A., Ryoo, M.S., Sindhwani, V., Lee, J., Vanhoucke, V., Florence, P.: Socratic models: Composing zero-shot multimodal reasoning with language. In: The Eleventh International Conference on Learning Representations (2023), <https://openreview.net/forum?id=G2Q2Mh3avow> [1](#)
53. Zhai, X., Mustafa, B., Kolesnikov, A., Beyer, L.: Sigmoid loss for language image pre-training. In: Proceedings of the IEEE/CVF International Conference on Computer Vision. IEEE (2023) [1](#), [4](#), [7](#)
54. Zhang, B., Li, K., Cheng, Z., Hu, Z., Yuan, Y., Chen, G., Leng, S., Jiang, Y., Zhang, H., Li, X., Jin, P., Zhang, W., Wang, F., Bing, L., Zhao, D.: Videollama 3: Frontier multimodal foundation models for image and video understanding. arXiv preprint arXiv:2501.13106 (2025), <https://arxiv.org/abs/2501.13106> [1](#), [5](#), [7](#), [8](#)
55. Zhang, Y., Wu, J., Li, W., Li, B., Ma, Z., Liu, Z., Li, C.: Video instruction tuning with synthetic data (2024), <https://arxiv.org/abs/2410.02713> [1](#)
56. Zhou, K., Yang, J., Loy, C.C., Liu, Z.: Conditional prompt learning for vision-language models. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (2022) [3](#), [4](#)
57. Zhu, B., Niu, Y., Han, Y., Wu, Y., Zhang, H.: Prompt-aligned gradient for prompt tuning. In: Proceedings of the IEEE/CVF International Conference on Computer Vision (2023) [3](#)
58. Zhu, D., Sun, Z., Li, Z., Shen, T., Yan, K., Ding, S., Kuang, K., Wu, C.: Model tailor: Mitigating catastrophic forgetting in multi-modal large language models. In: International Conference on Machine Learning (2024) [12](#), [13](#)
59. Zohar, O., Wang, X., Dubois, Y., Mehta, N., Xiao, T., Hansen-Estruch, P., Yu, L., Wang, X., Juefei-Xu, F., Zhang, N., Yeung-Levy, S., Xia, X.: Apollo: An exploration of video understanding in large multimodal models. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (2025) [4](#)