

4D-VLA: Spatiotemporal Vision-Language-Action Pretraining with Cross-Scene Calibration

Jiahui Zhang^{1*} Yurui Chen^{1*} Yueming Xu¹ Ze Huang¹ Yanpeng Zhou²
Yu-Jie Yuan² Xinyue Cai² Guowei Huang² Xingyue Quan² Hang Xu² Li Zhang^{1†}

¹ School of Data Science, Fudan University ² Huawei Noah's Ark Lab

<https://github.com/LogosRoboticsGroup/4D-VLA>

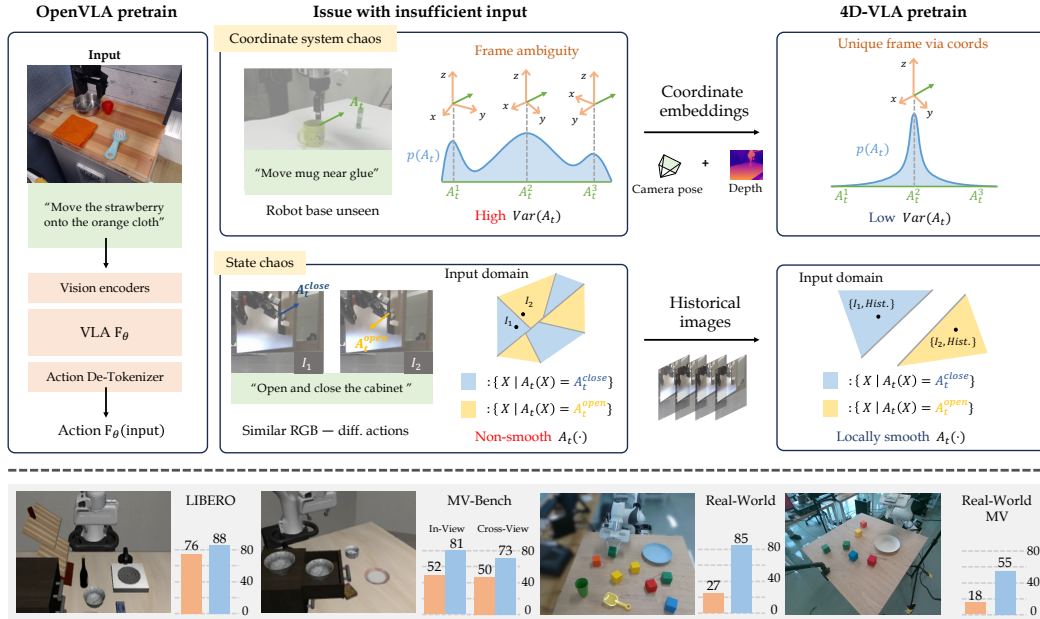


Figure 1: **Top:** Our pretraining design philosophy highlights that prior methods often lack key cues in their input for accurate action inference. This leads to target action distributions $A_t(\cdot)$ exhibiting high variance or non-smoothness, which negatively impacts pretraining performance. A rough analysis shows that in the DROID dataset, 67% of the samples have the robot's base occluded, causing coordinate system chaos. **Bottom:** We verify our method in both simulated and real-world robotic settings and report the performance for the **OpenVLA** baseline and our **4D-VLA** approach.

Abstract

Leveraging diverse robotic data for pretraining remains a critical challenge. Existing methods typically model the dataset's action distribution using simple observations as inputs. However, these inputs are often incomplete, resulting in a dispersed conditional action distribution—an issue we refer to as coordinate system chaos and state chaos. This inconsistency significantly hampers pretraining efficiency. To address this, we propose **4D-VLA**, a novel approach that effectively integrates 4D information into the input to mitigate these sources of chaos. Our model introduces

*Equal contribution. †Corresponding author (lizhangfd@fudan.edu.cn).

depth and temporal information into visual features with sequential RGB-D inputs, aligning the coordinate systems of the robot and the scene. This alignment endows the model with strong spatiotemporal reasoning capabilities while minimizing training overhead. Additionally, we introduce *memory bank sampling*, a frame sampling strategy designed to extract informative frames from historical images, further improving effectiveness and efficiency. Experimental results demonstrate that our pretraining method and architectural components substantially enhance model performance. In both simulated and real-world experiments, our model achieves a significant increase in success rate over OpenVLA [1]. To further assess spatial perception and generalization to novel views, we introduce *MV-Bench*, a multi-view simulation benchmark. Our model consistently outperforms existing methods, demonstrating stronger spatial understanding and adaptability.

1 Introduction

The emergence of pretrained vision-language models has established an effective framework for aligning human language with visual data, advancing embodied intelligence. Meanwhile, the open-sourcing of diverse robotic datasets [2, 3, 4, 5] has enabled data-driven robot control. However, efficiently extracting useful information from these datasets remains a challenge for improving generalization across diverse scenarios.

A real-world action distribution can be interpreted as a response function conditioned on observations or input, denoted as $A_t(\text{input})$. The objective of pretraining is to learn a model $F_\theta(\text{input})$ that approximates this function using large-scale data $A_t^{\text{data}}(\text{input})$. However, existing pretraining paradigms often suffer from incomplete or under-informative input, lacking critical contextual cues required for reliable action reasoning. Consequently, the resulting target distribution $A_t(\cdot)$ may exhibit undesirable characteristics—such as lack of smoothness, high variance, or multimodality—which hinder the model’s ability to learn robust and generalizable behaviors.

Previous approaches, such as OpenVLA [1], use only a single RGB image and a textual instruction as input. This limited input setting leads to two prominent types of chaos. The first is coordinate system chaos, which arises when actions are defined in the robot’s coordinate frame, yet the visual input lacks sufficient spatial context. For instance, if the image does not fully capture the robot’s body, it becomes challenging to infer the robot’s exact position and orientation. The second is state chaos, which arises in scenarios where a single frame lacks the necessary temporal or contextual cues to resolve action ambiguity. This includes symmetric trajectories—where it is difficult to infer the direction of motion—as well as cases where visually similar observations correspond to entirely different actions. These ambiguities hinder the effectiveness of pretraining, as illustrated in Fig. 1. In contrast, HPT [6] extends the input by incorporating some dataset-specific parameters, which helps address the issue of inconsistent coordinate systems across different datasets. However, this approach lacks scalability and increases the complexity of training.

To address this, we propose *4D-VLA*, a framework that integrates 4D spatiotemporal information to resolve such ambiguities. By combining spatial coordinate embeddings with a 3D-aware module, it generates spatial vision tokens that align the robot’s coordinate system with the scene, enhancing 3D perception. This also enables efficient encoding of multiple historical frames, improving temporal reasoning. Additionally, we introduce *memory bank sampling*, a frame sampling strategy that selects key frames based on historical similarity, boosting model efficiency.

To further investigate the spatial understanding and generalization of VLA models, we introduce *MV-Bench*, a multi-view dataset that evaluates performance across diverse viewpoints. Our approach enables robust pretraining, improving generalization to novel scenarios while outperforming baselines.

Our contributions are: (i) We propose *4D-VLA*, an efficient VLA model that integrates a spatial module with vision features to generate 3D-aware spatial vision tokens, effectively mitigating coordinate system and state chaos, thereby significantly enhancing pretraining efficiency. Additionally, we introduce *memory bank sampling*, a simple but effective method for historical information sampling. (ii) Our model has been validated both in the simulated and real-world environment, demonstrating its superiority. (iii) We develop a multi-view simulation dataset *MV-Bench* to evaluate spatial understanding and generalization, on which our model achieves outstanding performance.

2 Related works

Vision-language models Recent advancements in vision-language models (VLMs) have significantly enhanced the integration of vision and language understanding across diverse domains. Models like Flamingo [7], LLaVA [8], and BLIP-2 [9] focus on aligning text and image feature spaces to facilitate effective image understanding. More recently, there has been growing interest in expanding language models to support multi-image inputs, enabling them to handle more complex tasks and real-world scenarios. For instance, [10, 11, 12, 13] utilize multi-image sequences to capture temporal and action-related dynamics across frames, providing robust video comprehension. In addition, some recent models have begun incorporating 3D information to bolster spatial reasoning, as seen in 3D-LLM [14] and Scene-LLM [15]. These models leverage 3D inputs to enhance understanding of spatial relationships within a scene. However, none of these models explicitly integrate 4D spatiotemporal information to fully capture both spatial and temporal dynamics within the architecture.

Vision-language-action models The vision-language-action (VLA) model represents a significant advancement in vision-language research, enabling more complex and interactive tasks aimed at facilitating real-world environment interactions. [16, 17, 18] complete tasks by directly predicting trajectories. [19, 20, 21, 22, 23, 24] predict the robot’s current actions to enable closed-loop control, while [25, 26] enhance action prediction by training a world model to forecast future states. Some works [27, 28] also drives VLA policies using simple concatenation of historical observations.

Recent works leverage diverse robotic datasets from various scenes and robot types to pretrain models for better generalization in novel environments. [29, 30, 3, 1] use the single image as input and are pretrained on large-scale datasets, while Octo [31] incorporates historical context and uses a diffusion head to predict the next n actions. HPT [6] addresses the heterogeneity among different datasets by introducing dataset-specific parameters during pretraining to improve training efficiency. However, these methods overlook that the inefficiency in prior pretraining arises from insufficient input context, resulting in a high variance of the conditioned action distribution $A_t(\cdot)$ and ultimately hindering pretraining effectiveness. Our approach tackles this issue by introducing 4D information to mitigate coordinate system chaos and state chaos. This enables the model to learn meaningful action distributions from diverse datasets, thereby enhancing performance.

3 Method

This section provides a comprehensive overview of our proposed *4D-VLA*. As shown in Fig. 2, our model processes sequential RGB-D images as input, converting them into corresponding spatial vision tokens. These tokens, together with task-specific text tokens, serve as feature inputs for the subsequent Transformer. After decoding through the VLM Transformer and action head, it ultimately generates the action output.

3.1 Preliminary

Problem definition The vision-language action (VLA) model takes a language instruction as input and aims to control a robot to accomplish the specified task. Specifically, VLA with a low-level control policy refers to a class of models that use the current observations as input to predict an action for the robot in its present state, enabling end-to-end, closed-loop control. The action is defined by three components: $\Delta \mathbf{x} \in \mathbb{R}^3$, $\Delta \boldsymbol{\theta} \in \mathbb{R}^3$, and $g \in [0, 1]$, representing the control translation, rotation offset, and the open-close state of the robot’s end-effector, respectively.

Vision-language model backbone We leverage a pretrained large vision-language model (VLM) as the backbone, specifically InternVL-4B [12], which consists of a text tokenizer $\mathcal{T}$, a vision encoder $\mathcal{E}$, and a Transformer decoder $\mathcal{D}$. The vision encoder processes visual observations, which are subsequently compressed by an MLP projector $\mathcal{P}$ to generate vision embeddings, while text inputs are tokenized and embedded to form structured textual tokens. These multimodal tokens are then fed into the decoder $\mathcal{D}$ for next-token prediction. This backbone provides a robust foundation for aligning visual information with the shared semantic space, enabling effective robot action generation.

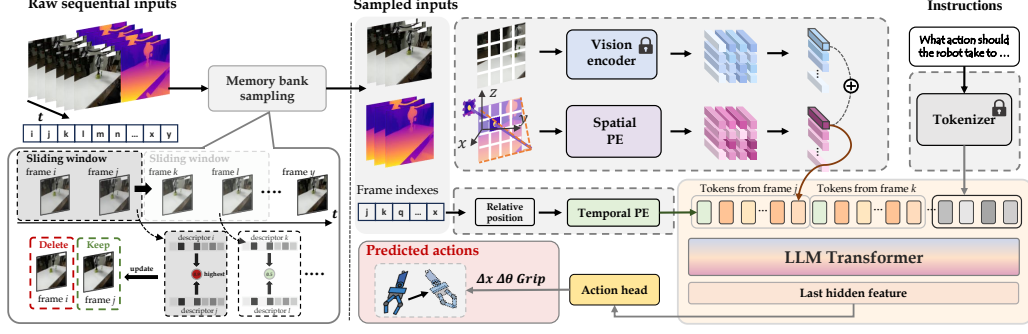


Figure 2: **Our 4D-VLA pipeline.** Our *memory bank sampling* method selects informative frames from sequential RGB-D inputs. A vision encoder with 3D coordinate embeddings generates spatial-aware tokens, which are fused into a 4D spatiotemporal representation. Combined with text tokens, these are processed by the LLM to decode actions via an action head.

3.2 Spatial-aware visual tokens

A reasonable action prediction requires awareness of both semantic perception and spatial perception of the scene. A fundamental aspect of spatial perception is that the robot must accurately determine its relative position within the scene, effectively resolving coordinate system chaos. To address this issue while enhancing spatial perception without compromising semantic awareness, we fuse the coordinate information derived from depth with vision tokens, resulting in spatial vision tokens.

In our method, the input image $\mathbf{I} \in \mathbb{R}^{3 \times h \times w}$ is first encoded by $\mathcal{E}$ into a feature map with a downsampling rate of c , yielding $\mathbf{f}_v = \mathcal{E}(\mathbf{I}) \in \mathbb{R}^{k \times \frac{h}{c} \times \frac{w}{c}}$, where k denotes the feature dimension. Next, we obtain a downsampled depth map $\mathbf{D} \in \mathbb{R}^{\frac{h}{c} \times \frac{w}{c}}$, which assigns depth values to each corresponding feature volume. Using the camera’s extrinsic $[\mathbf{R$

The algorithm is detailed in Alg. 1, which sequentially traverses image groups while maintaining a similarity queue, ensuring each newly added frame has lower similarity than the current maximum.

Temporal positional encoding. Since *memory bank sampling* follows a non-uniform strategy, it is essential to encode the temporal position of each sampled spatial vision token relative to the current frame. To enhance the model’s flexibility and generalization, we introduce a time encoding token e^T , which captures the relative temporal offset between the historical and current frames, $e_j^T = \mathcal{E}_T(t - j)$, where j denotes the timestamp of a historical frame, t represents the current observation time, and $\mathcal{E}_t$ is the learnable temporal encoding function. Accordingly, the final input token set is structured as:

$$\mathcal{X} = \bigcup_{i \in \mathcal{H}} [e_i^T \mid e_i^{ST}] \cup \{e^{text}\}, \quad (2)$$

where e^{text} represents the text instruction tokens. Each sampled frame i contributes a token pair $[e_i^T \mid e_i^{ST}]$, with the temporal encoding token preceding the spatial vision token. This design ensures that the model effectively captures temporal relationships while maintaining spatial-awareness.

3.4 Loss functions

To accelerate control policy generation, we use an action head with two MLP layers that predict the action $[\Delta\hat{x}, \Delta\hat{\theta}, \hat{g}]$ using the hidden features of the VLM Transformer’s last token.

Our total training loss can be written as follows:

$$\mathcal{L} = \mathcal{L}_t + \mathcal{L}_r + \mathcal{L}_g + \lambda_d \mathcal{L}_d, \quad (3)$$

where the translation loss $\mathcal{L}_t = \|\Delta\hat{x} - \Delta x\|_2$, the rotation loss $\mathcal{L}_r = \|\Delta\hat{\theta} - \Delta\theta\|_2$

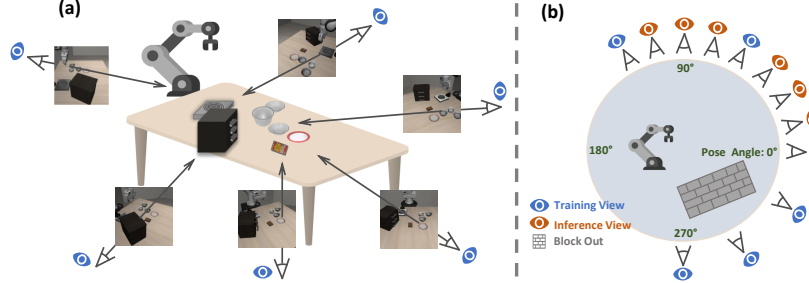


Figure 3: **Our MV-Bench camera setting.** We select 6 diverse viewpoints as training views and render images for all LIBERO-SPATIAL tasks. Novel inference views are placed near the training views. To avoid occlusion from the black box, test views in blocked areas are excluded.

Method	Spatial	Object	Goal	Long	Avg.
UniAct-0.5B [33] [†]	64.5	77.5	68.0	46.5	64.1
SparseVLM [34] [†]	79.8	67.0	72.6	39.4	64.7
FastV [35] [†]	83.4	84.0	74.2	51.6	73.3
VLA-Cache [36] [†]	83.8	85.8	76.4	52.8	74.7
DiffusionPolicy [16]	78.3 ± 1.1	92.5 ± 0.7	68.3 ± 1.2	50.5 ± 1.3	72.4 ± 0.7
TraceVLA [37]	84.6 ± 0.2	85.2 ± 0.4	75.1 ± 0.3	54.1 ± 1.0	74.8 ± 0.5
SpatialVLA [38]	88.2 ± 0				

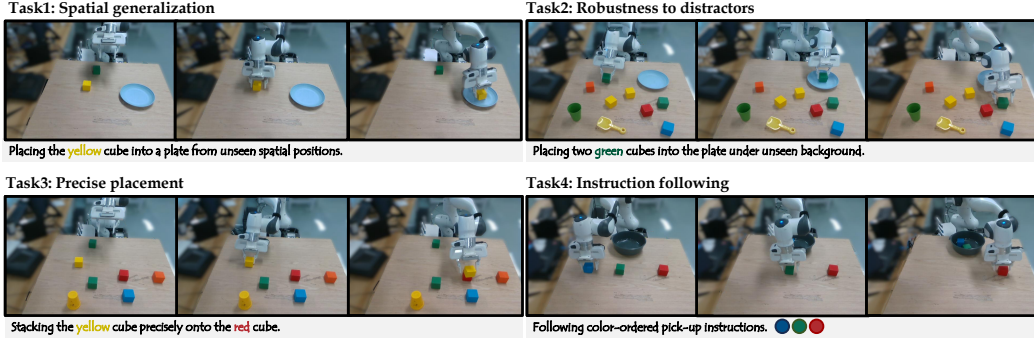


Figure 4: **Our real-world experiment settings.** These settings aim to evaluate the model’s spatial generalization, robustness to distractors, precision in placement, and ability to follow instructions. Each row presents a 3-frame execution snapshot.

4.3 LIBERO evaluation

After pretraining, we fine-tune and conduct close-loop testing in the LIBERO simulation environment, using task success rate as the metric to evaluate the performance.

Evaluation protocol. For each task in LIBERO, there are 10 test subtasks, and each subtask contains 50 different object layouts. This requires 1500 simulation tracks to evaluate a single task. The random seed is used to alter the initial state of the objects. To evaluate each task, we randomly sample 3 different seeds to calculate the mean and standard deviation of the success rate.

Fine-tuning details. Unlike the pretraining phase, we used the simplest input settings to enable our model to learn the interaction effects between 3D information and historical data on actions. In the subsequent fine-tuning phase, we set the number of sequential frames $k = 5$ with a window size $n = 20$. The current state RGB-D image is resized to 448×448 , while historical images are resized to 224×224 to optimize memory usage and model efficiency. Following [1], we normalize the ground truth action label using the 99th and 1st percentiles. We employed a cosine learning rate scheduler with a learning rate of $4e-5$, using a batch size of 128 and training for 20 epochs. λ_d is set to 0. During training, we activated all network parameters except the vision encoder.

Experimental results. As shown in Tab. 1, our <

Method	Task 1	Task 2	Task 3	Task 4	Avg.
OpenVLA [1]	45.00	22.50	30.00	13.33	27.70
Base VLA	35.00	20.00	5.00	2.67	15.67
+ Pretraining	60.00	60.00	40.00	28.00	47.00
+ Pretraining + Coord.	75.00	60.00	85.00	34.67	63.67
+ Pretraining + Hist.	80.00	77.50	70.00	36.00	65.88
+ Pretraining + Coord. + Hist. (<i>full model</i>)	90.00	82.50	90.00	80.00	85.63

Table 3: **Real-world evaluation results.** We incrementally improve the Base VLA by adding pretraining, coordinate encoding, and historical frames selected via *memory bank sampling*.

training and trained the model for 20 epochs. For Task 4, involving color-ordered pick-up instructions, we used 5 color sequences with 10 demonstrations each, totaling 50 trajectories.

Task descriptions. As shown in Fig. 4, we design 4 real-world manipulation tasks to evaluate different aspects of the model’s spatial reasoning and generalization capabilities. **Task 1: Spatial generalization.** The robot is tasked with placing a yellow cube into a plate from positions not seen during training, evaluating its ability to generalize across novel spatial positions. **Task 2: Robustness to distractors.** The robot is asked to place two green cubes into the plate in scenes with cluttered backgrounds. This task evaluates the model’s robustness to environments distractors. **Task 3: Precise placement.**

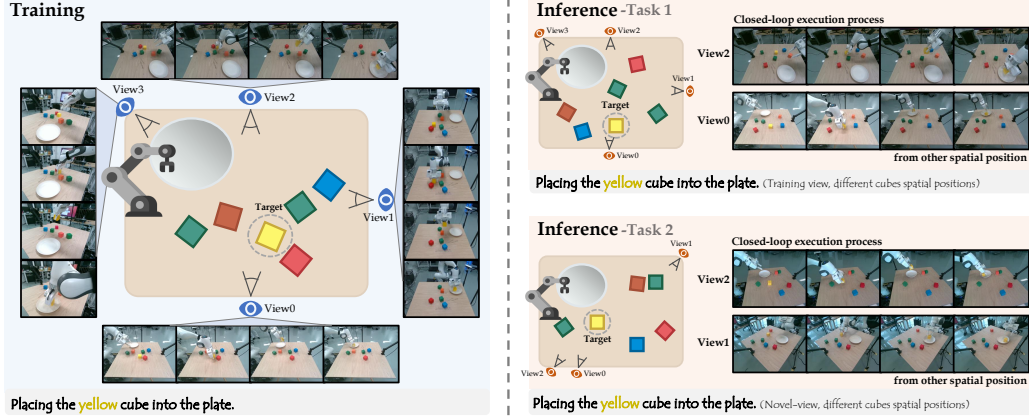


Figure 5: **Our multi-view real-world experiment settings.** These settings aim to evaluate the model’s out-of-distribution and novel-view generalization ability.

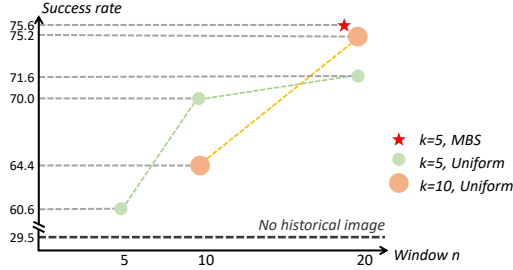
Method	In-View				Cross-View			Avg.
<i>Angle</i>	0°	90°	180°	225°	$\Delta 15^\circ (-15^\circ)$	$\Delta 25^\circ (-25^\circ)$	$\Delta 45^\circ (135^\circ)$	
OpenVLA [1]	25	15	30	10	30	10	5	18
4D-VLA (Ours)	60	50	65	65	50	55	40	55

Table 4: **Real-world multi-view evaluation.** We test our model’s spatial generalization across varying viewpoints and object layouts. 4

performance of the model? (iii) **The impact of coordinate system chaos** -How does coordinate system inconsistency affect the model’s performance, and can the introduction of a 3D representation mitigate this issue? The first question is addressed in Sec. 5.1; the second is covered in Appx. 7.2, and the third in Appx. 7.3. Our discussion begins with the simple model, i.e., InternVL-4B with an MLP action head, using single RGB image as the vision input.

5.1 Exploring historical information utilization

As highlighted in Octo [31], incorporating historical context can significantly boost model performance. However, this has been underexplored. We conduct a comprehensive analysis to assess its true impact on model effectiveness. We define window size n as the number of historical frames available to the model and k as the number of frames sampled from them. To systematically investigate their impact, we design two experimental settings: (i) fixing n while progressively increasing k , and (ii) fixing k while expanding n .



References

- [1] Moo Jin Kim, Karl Pertsch, Siddharth Karamcheti, Ted Xiao, Ashwin Balakrishna, Suraj Nair, Rafael Rafailov, Ethan Foster, Grace Lam, Pannag Sanketi, et al. Openvla: An open-source vision-language-action model. *arXiv preprint*, 2024. 2, 3, 6, 7, 8, 9
- [2] Alexander Khazatsky, Karl Pertsch, Suraj Nair, Ashwin Balakrishna, Sudeep Dasari, Siddharth Karamcheti, Soroush Nasiriany, Mohan Kumar Srirama, Lawrence Yunliang Chen, Kirsty Ellis, Peter David Fagan, Joey Hejna, Masha Itkina, Marion Lepert, Yecheng Jason Ma, Patrick Tree Miller, Jimmy Wu, Suneel Belkhale, Shivin Dass, Huy Ha, Arhan Jain, Abraham Lee, Youngwoon Lee, Marius Memmel, Sungjae Park, Ilija Radosavovic, Kaiyuan Wang, Albert Zhan, Kevin Black, Cheng Chi, Kyle Beltran Hatch, Shan Lin, Jingpei Lu, Jean Mercat, Abdul Rehman, Pannag R Sanketi, Archit Sharma, Cody Simpson, Quan Vuong, Homer Rich Walke, Blake Wulfe, Ted Xiao, Jonathan Heewon Yang, Arefeh Yavary, Tony Z. Zhao, Christopher Agia, Rohan Baijal, Mateo Guaman Castro, Daphne Chen, Qiuyu Chen, Trinity Chung, Jaimyn Drake, Ethan Paul Foster, Jensen Gao, David Antonio Herrera, Minho Heo, Kyle Hsu, Jiaheng Hu, Donovan Jackson, Charlotte Le, Yunshuang Li, Kevin Lin, Roy Lin, Zehan Ma, Abhiram Maddukuri, Suvir Mirchandani, Daniel Morton, Tony Nguyen, Abigail O’Neill, Rosario Scalise, Derick Seale, Victor Son, Stephen Tian, Emi Tran, Andrew E. Wang, Yilin Wu, Annie Xie, Jingyun Yang, Patrick Yin, Yunchu Zhang, Osbert Bast

- [5] Oier Mees, Lukas Hermann, Erick Rosete-Beas, and Wolfram Burgard. Calvin: A benchmark for language-conditioned policy learning for long-horizon robot manipulation tasks. In *RA-L*, 2022. 2
- [6] Lirui Wang, Xinlei Chen, Jialiang Zhao, and Kaiming He. Scaling proprioceptive-visual learning with heterogeneous pre-trained transformers. *arXiv preprint*, 2024. 2, 3
- [7] Jean-Baptiste Alayrac, Jeff Donahue, Pauline Luc, Antoine Miech, Iain Barr, Yana Hasson, Karel Lenc, Arthur Mensch, Katherine Millican, Malcolm Reynolds, et al. Flamingo: a visual language model for few-shot learning. In *NeurIPS*, 2022. 3
- [8] Haotian Liu, Chunyuan Li, Qingyang Wu, and Yong Jae Lee. Visual instruction tuning. In *NeurIPS*, 2023. 3
- [9] Junnan Li, Dongxu Li, Silvio Savarese, and Steven C. H. Hoi. BLIP-2: Bootstrapping language-image pre-training with frozen image encoders and large language models. In *ICML*, 2023. 3
- [10] Tim Dettmers, Mike Lewis, Younes Belkada, and Luke Zettlemoyer. Gpt3. int8 (): 8-bit matrix multiplication for transformers at scale. In *NeurIPS*, 2022. 3
- [11] Jinze Bai, Shuai Bai, Shusheng Yang, Shijie Wang, Sinan Tan, Peng Wang, Junyang Lin, Chang Zhou, and Jingren Zhou. Qwen-vl: A frontier large vision-language model with versatile abilities. *arXiv preprint*, 2023. 3
- [12] Zhe Chen, Jiannan Wu, Wenhai Wang, Weijie Su, Guo Chen, Sen Xing, Muyan Zhong, Qinglong Zhang, Xizhou Zhu, Lewei Lu, et al. Internvl: Scaling up vision foundation models and aligning for generic visual-linguistic tasks. In *CVPR*, 2024. 3, 4
- [13] Guo Chen, Yin-Dong Zheng, Jiahao Wang, Jilan Xu, Yifei Huang, Junting Pan, Yi Wang, Yali Wang, Yu Qiao, Tong Lu, et al. Videollm: Modeling video sequence with large language models. *arXiv preprint*, 2023. 3
- [14] Yining Hong, Haoyu Zhen, Peihao Chen, Shuhong Zheng, Yilun Du, Zhenfang Chen, and Chuhan Gan. 3d-llm: Injecting the 3d world into large language models. In *NeurIPS*

- [23] Mengda Xu, Zhenjia Xu, Yinghao Xu, Cheng Chi, Gordon Wetzstein, Manuela Veloso, and Shuran Song. Flow as the cross-domain manipulation interface. *arXiv preprint*, 2024. 3
- [24] Jiangyong Huang, Silong Yong, Xiaojian Ma, Xiongkun Linghu, Puhao Li, Yan Wang, Qing Li, Song-Chun Zhu, Baoxiong Jia, and Siyuan Huang. An embodied generalist agent in 3d world. In *ICML*, 2024. 3
- [25] Haoyu Zhen, Xiaowen Qiu, Peihao Chen, Jincheng Yang, Xin Yan, Yilun Du, Yining Hong, and Chuang Gan. 3d-vla: A 3d vision-language-action generative world model. In *ICML*, 2024. 3
- [26] Hongtao Wu, Ya Jing, Chilam Cheang, Guangzeng Chen, Jiafeng Xu, Xinghang Li, Minghuan Liu, Hang Li, and Tao Kong. Unleashing large-scale video generative pre-training for visual robot manipulation. In *ICLR*, 2024. 3
- [27] Ilija Radosavovic, Bike Zhang, Baifeng Shi, Jathushan Rajasegaran, Sarthak Kamat, Trevor Darrell, Koushil Sreenath, and Jitendra Malik. Humanoid locomotion as next token prediction. *Advances in neural information processing systems*, 37:79307–79324, 2024. 3
- [28] Ilija Radosavovic, Baifeng Shi, Letian Fu, Ken Goldberg, Trevor Darrell, and Jitendra Malik. Robot learning with sensorimotor pre-training. In *CoRL*, 2023. 3
- [29] Anthony Brohan, Noah Brown, Justice Carbajal, Yevgen Chebotar, Joseph Dabis, Chelsea Finn, Keerthana Gopalakrishnan, Karol Hausman, Alex Herzog, Jasmine Hsu, Julian Ibarz, Brian Ichter, Alex Irpan, Tomas Jackson, Sally Jesmonth, Nikhil Joshi, Ryan Julian, Dmitry Kalashnikov, Yuheng Kuang, Isabel Leal, Kuang-Huei Lee, Sergey Levine, Yao Lu, Utsav Malla, Deeksha Manjunath, Igor Mordatch, Ofir Nachum, Carolina Parada, Jodilyn Peralta, Emily Perez, Karl Pertsch, Jornell Quiambao, Kanishka Rao, Michael Ryoo, Grecia Salazar, Pannag Sanketi, Kevin Sayed, Jaspiar

- [36] Siyu Xu, Yunke Wang, Chenghao Xia, Dihao Zhu, Tao Huang, and Chang Xu. Vla-cache: Towards efficient vision-language-action model via adaptive token caching in robotic manipulation. *arXiv preprint*, 2025. 6
- [37] Ruijie Zheng, Yongyuan Liang, Shuaiyi Huang, Jianfeng Gao, Hal Daumé III, Andrey Kolobov, Furong Huang, and Jianwei Yang. Tracevla: Visual trace prompting enhances spatial-temporal awareness for generalist robotic policies. *arXiv preprint*, 2024. 6
- [38] Delin Qu, Haoming Song, Qizhi Chen, Yuanqi Yao, Xinyi Ye, Yan Ding, Zhigang Wang, JiaYuan Gu, Bin Zhao, Dong Wang, et al. Spatialvla: Exploring spatial representations for visual-language-action model. *arXiv preprint*, 2025. 6
- [39] Stephen James, Kentaro Wada, Tristan Laidlow, and Andrew J Davison. Coarse-to-fine q-attention: Efficient learning for visual robotic manipulation via discretisation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 13739–13748, 2022. 22
- [40] Dantong Niu, Yuvan Sharma, Giscard Biamby, Jerome Quenum, Yutong Bai, Baifeng Shi, Trevor Darrell, and Roei Herzig. Llarva: Vision-action instruction tuning enhances robot learning. In *Conference on Robot Learning*, pages 3333–3355. PMLR, 2025. 22
- [41] Mohit Shridhar, Lucas Manuelli, and Dieter Fox. Perceiver-actor: A multi-task transformer for robotic manipulation. In *Conference on Robot Learning*, pages 785–799. PMLR, 2023. 22
- [42] Xiaoqian Shen, Yunyang Xiong, Changsheng Zhao, Lemeng Wu, Jun Chen, Chenchen Zhu, Zechun Liu, Fanyi Xiao, Balakrishnan Varadarajan, Florian Bordes, et al. Longvu: Spatiotemporal adaptive compression for long video-language understanding. In *Forty-second International Conference on Machine Learning*. 22, 23
- [43] Yanwei Li, Chengyao Wang, and Jiaya Jia. Llama-vid: An image is worth 2 tokens in large language models. In *European Conference on Computer Vision*, pages 323–340. Springer, 2024. 22, 23

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: We have clearly stated the contribution and scope of the paper in the abstract.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: We have discussed the limitations of the work in the conclusion part.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach

Justification: We have provided the theoretical assumption and proof of this paper in the method part.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [\[Yes\]](#)

Justification: We have disclosed all the information needed to reproduce the main experimental results of the paper in the experiment part.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [No]

Justification: We have not released the code of this paper in the submission.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: We have specified all the training and test details of this paper in the experiment part.

Guidelines:

- The answer NA

- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: We have provided sufficient information on the computer resources of this method in the experiment part.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: This paper has not released data or models that have a high risk for misuse.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. **New assets**

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: Our paper does not release new assets, including data and models. We also did not publish or submit the code.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. **Crowdsourcing and research with human subjects**

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: The research topic of this paper is vision language action model, not involving crowdsourcing nor research with human subjects.

Guid

- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. **Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification: LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.

7 Appendix section

7.1 Comparison with ARM4R

Our approach differs from ARM4R in both design and practice. Instead of representation learning with full 3D point clouds and proprioception, we adopt an *end-to-end, action-centric* pretraining pipeline that maps vision–language inputs directly to low-level actions. Inputs are RGB with depth-aligned spatial coordinates (2D patches + (x, y, z) embeddings), keeping the model lightweight and VLM-compatible while avoiding point-cloud encoders.

Method	Meet-off-grill	Push-buttons	Place-wine	Open-drawer	Avg.
C2FARM-BC [39]	20.0	72.0	18.0	20.0	32.5
LLARVA [40]	80.0	56.0	12.0	60.0	52.0
PerAct [41]	84.0	48.0	12.0	80.0	56.0
ARM4R [32]	94.4	67.2	36.0	88.8	71.9
4D-VLA(Ours)	95.2	79.2	45.6	89.6	77.4

Table 6: **RLBench results under the ARM4R protocol.**

Analysis. To balance efficiency and spatial coverage, we use our streaming *memory bank sampling* strategy that keeps 5 history frames within a 20-frame context. The current step uses dual views (front + wrist), while history stores only the front view. We also adopt OpenVLA’s random-crop augmentation. This configuration provides strong spatial grounding and good runtime

Memory bank sampling. We compare our *memory bank sampling* with three video-style samplers: *Adaptive Pooling* [42], *Grid Pooling* [43], and a *per frame Q-Former*, using InternVL2.5-4B on Libero-Spatial (20 epochs). We report success, latency, and peak memory in Tab. 8. *Adaptive Pooling* caches the full trajectory and uses two thresholds (0.99 for frame selection and 0.96 for spatial compression). We omit its second-stage cross-modal filter due to low language diversity (10 tasks) and observe hyperparameter sensitivity and variable-length tokens that increase memory and reduce accuracy. *Grid Pooling* replaces language tokens with 8 learnable queries and compresses each frame to a 2×2 grid (4 tokens), yielding $20 \times (4+8)$ historical tokens; it is simple but nonadaptive and can drop key spatial cues. *Q-Former (per frame)* uses 10 learnable queries per frame to extract features but converges slowly and discards too much spatial context, which hurts success. *MBS (ours)* operates online with a window of 20 and a memory of 5. It selects informative key frames from the stream, stores only 5 history frames at 224×224 , and keeps the current frame at 448×448 . As a result, MBS achieves the highest success (0.866) with competitive cost (160.0 ms and 8,682.9 MB). Single frame is faster (76.5 ms) but weaker (0.738). The other baselines are slower and less accurate. MBS fits VLA’s streaming and causal nature because it avoids full-clip caching, uses relative temporal encoding, and remains compatible with closed-loop control. It yields stronger long horizon performance without extra retraining or heavy compute.

7.3 Coordinate system chaos impact analysis

To assess the impact of coordinate system chaos on VLA’s performance, we conduct a controlled experiment. The experiment consists of two main steps: First, we deliberately introduce controlled chaos into the LIBERO environment. Then, we compare the performance of models with and without 3D information. Our findings demonstrate that chaotic coordinate transformations significantly degrade model performance while incorporating 3D information effectively alleviates it.

Chaos generation. To simulate the diverse viewpoints in the pretraining dataset—where the robot’s body is absent from the image, and

Experimental results. We control chaos levels by adjusting the magnitude of random rotations. Level 0 applies no rotation, while levels 1–3 introduce random z-axis rotations of 15° , 30° , and 90° , respectively. Translation t is randomly set within a range of 0.5. As shown in Fig. 7, without chaos, both models perform well, with 3D information further boosting success rates. Notably, the 3D model shows lower variance across viewpoints. As chaos increases, the non-3D model’s performance drops sharply, while the 3D model maintains relatively high success—highlighting the value of 3D cues in handling coordinate system chaos.