



WorldVLA: Towards Autoregressive Action World Model

Jun Cen^{1,2,3}, Chaohui Yu¹, Hangjie Yuan^{1,2,3}, Yuming Jiang¹, Siteng Huang^{1,2}, Jiayan Guo¹, Xin Li^{1,2}, Yibing Song¹, Hao Luo^{1,2}, Fan Wang¹, Deli Zhao^{1,2}, Hao Chen³

¹DAMO Academy, Alibaba Group ²Hupan Lab ³Zhejiang University

We present WorldVLA, an autoregressive action world model that unifies action and image understanding and generation. Our WorldVLA integrates Vision-Language-Action (VLA) model and world model in one single framework. The world model predicts future images by leveraging both action and image understanding, with the purpose of learning the underlying physics of the environment to improve action generation. Meanwhile, the action model generates the subsequent actions based on image observations, aiding in visual understanding and in turn helps visual generation of the world model. We demonstrate that WorldVLA outperforms standalone action and world models, highlighting the mutual enhancement between the world model and the action model. In addition, we find that the performance of the action model deteriorates when generating sequences of actions in an autoregressive manner. This phenomenon can be attributed to the model’s limited generalization capability for action prediction, leading to the propagation of errors from earlier actions to subsequent ones. To address this issue, we propose an attention mask strategy that selectively masks prior actions during the generation of the current action, which shows significant performance improvement in the action chunk generation task.

Date: June 27, 2025

Code: <https://github.com/alibaba-damo-academy/WorldVLA>

Correspondence: cenjun.cen@alibaba-inc.com



1 Introduction

The development of Vision-Language-Action (VLA) models has emerged as a significant focus within robotics action model research (Brohan et al., 2023; Kim et al., 2024; Black et al., 2024). These models are constructed by augmenting large-scale pre-trained Multimodal Large Language Models (MLLMs) (Liu et al., 2023b; Li et al., 2024; Zhang et al., 2025; Bai et al., 2025) with either an action head or additional action expert module to generate actions. MLLMs contribute robust capabilities in perception and decision making, enabling VLA models to exhibit enhanced generalization across a wide range of robotic tasks (Black et al., 2024; Intelligence et al., 2025). Nevertheless, a notable limitation persists: these models often lack a comprehensive understanding of actions, as actions are treated solely as outputs but not being integrated as inputs for deeper analysis. In contrast, world models demonstrate the ability to predict future visual states based on current observations and actions, thereby achieving a dual understanding of both visual information and behavioral dynamics (Ha and Schmidhuber, 2018; Agarwal et al., 2025; Wu et al., 2025). Despite this advantage, world models are constrained by their inability to directly generate action outputs, resulting in a functional gap that limits their application in scenarios requiring explicit action planning.

To address the constraints inherent in both Vision-Language-Action (VLA) models and world models, we introduce WorldVLA, an autoregressive action world model for unified action and image understanding and generation. As depicted in Fig. 1, WorldVLA employs three separate tokenizers to encode images, text, and actions. The tokens from different modalities are set to share the same vocabulary so that understanding and generation across these modalities can be unified within a single LLM architecture. The world model component captures the underlying physical dynamics of the environment by generating visual representations based on input actions. This process of action interpretation and environmental physics learning is essential for enabling effective decision making within the action model. Concurrently, the action model embedded within

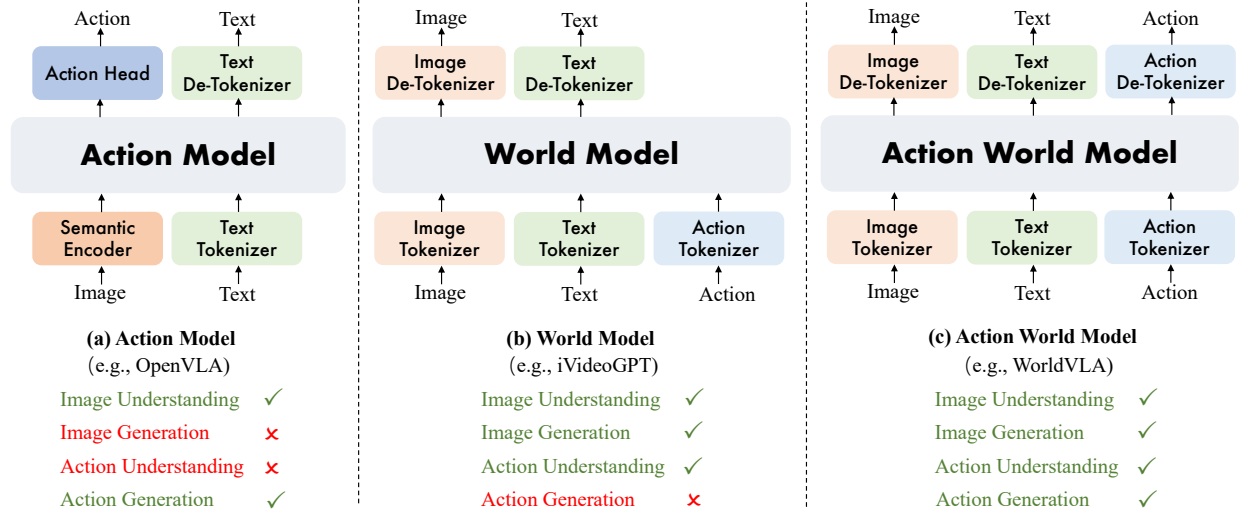


Figure 1 (a) Action model generates actions based on image understanding; (b) World model generates the image based on image and action understanding; (c) Action World Model unifies both image and action understanding and generation.

WorldVLA refines the understanding of visual data, thereby improving the precision of image generation performed by the world model. This bidirectional enhancement creates a more robust and comprehensive model capable of both understanding and generating actions and images.

Action chunking and parallel decoding have been demonstrated to significantly influence the performance of action models (Kim et al., 2025). However, we find that generating multiple actions in sequence leads to performance drop in autoregressive models. The primary reason for this is that pretrained multimodal language models have predominantly been exposed to images and text rather than actions, resulting in limited action generalization capabilities. In autoregressive models where subsequent actions are conditioned on preceding ones, error propagation becomes a critical issue, as the earlier incorrect predictions influence subsequent actions over time. To alleviate this issue, we propose an action attention masking strategy that selectively masks prior actions during the generation of current actions. This approach effectively mitigates error accumulation and yields substantial improvements in the task of action chunk generation.

The experiments on LIBERO benchmark show that our WorldVLA outperforms the action model with the same backbone by 4% grasping success rate. Further, compared to vanilla world model, our WorldVLA shows superior video generation capability and reduces Fréchet Video Distance (FVD) on LIBERO dataset by 10%. These results underscore the mutual benefits derived from integrating world and action models, highlighting the advantages of a unified framework for image and action comprehension and generation. In the context of action chunk generation, the grasping success rate decreases by 10% to 50% when employing a conventional autoregressive approach. However, the implementation of our attention masking strategy significantly mitigates this decrease, yielding a 4% to 23% improvement in grasping success rate.

In summary, our contributions are as follows:

- We propose WorldVLA, an autoregressive action world model that unifies action and image understanding and generation.
- We introduce an action attention masking strategy for the action chunk generation task in autoregressive models, addressing the challenge of action error accumulation when generating multiple actions in sequence.
- Our experiments demonstrate that WorldVLA outperforms the standalone action and world models, highlighting the mutual enhancement between the world model and action model. Additionally, the action attention masking strategy solves the performance degradation when generating action chunks and significantly improves grasping performance.

Table 1 Comparason of different action and video generative models. T: Text; V: Video; A: Action.

Model Type	Discrete	Continous	Input	Output
Action Model	OpenVLA (Kim et al., 2024)	π_0 (Black et al., 2024)	T + V	A
Video Prediction Model	MAGVIT (Yu et al., 2023)	SVD (Blattmann et al., 2023)	T + V	V
World Model	iVideoGPT (Wu et al., 2025)	DWS (He et al., 2025)	T + V + A	V
Action World Model	WorldVLA (ours)	UVA (Li et al., 2025)	T + V + A	V + A

2 Related Works

Our proposed WorldVLA is related to the action model, video prediction model, and world model. The difference between them are summarized in Table 1.

Vision-Language-Action Model. Behavior cloning (Pomerleau, 1988) is a classic imitation learning approach for robot manipulation, which learns a policy by mimicking expert observation-action pairs. Conventional architectures typically combine a vision backbone, such as ResNet (He et al., 2016) or Vision Transformer (Dosovitskiy et al., 2020), with an action head. The action head may consist of multilayer perceptrons (MLPs) (Rumelhart et al., 1986), query-based transformer decoders (Zhao et al., 2023), or diffusion-based policy heads (Chi et al., 2023). Recently, Vision-Language-Action (VLA) models have been proposed, utilizing pre-trained multi-modality large language models (MLLM) as the backbone (Brohan et al., 2022, 2023; Li et al., 2023; Huang et al., 2023; Belkhale and Sadigh, 2024; Wen et al., 2025; Zhen et al., 2024). These frameworks are equipped with either discrete action decoders (Kim et al., 2024; Pertsch et al., 2025) or continuous diffusion policy heads (Black et al., 2024; Wen et al., 2024) to predict actions. The internet-scale prior knowledge in MLLM enables effective generalization to unseen scenarios tasks for VLA models. Our proposed WorldVLA advances this paradigm by jointly generating actions and predicting future video frames, providing a comprehensive solution for understanding and generation.

Video Generation. Video generation plays a dual role in robotics. On one hand, some policy models generate the future video first and then generate the corresponding actions based on the generated video (Du et al., 2023; Ajay et al., 2023; Bu et al., 2024). Large-scale video data could be used for pre-training the future video generation part, as seen in approaches (Wu et al., 2023; Cheang et al., 2024). Here, video generation serves as a mechanism for visual imagination and planning, providing valuable insights that improve downstream policy generation (Cen et al., 2024). On the other hand, video generation models can act as world models, simulating diverse future scenarios (Ha and Schmidhuber, 2018). Such world models are widely utilized to generate varied training data (Agarwal et al., 2025), support model-based reinforcement learning algorithms (Wu et al., 2025), and aid in selecting the most suitable policies from a pool of generated options (Li et al., 2025; Bar et al., 2024). In this work, we show that our WorldVLA enables precise control over video generation through action inputs, while also demonstrating that video generation significantly enhances the quality of action generation.

Unified Understanding and Generation Model. Most multi-modality large language models (MLLMs) are designed to perform visual understanding tasks, where the model generates textual responses based on combined image and language inputs (Liu et al., 2023b; Li et al., 2024; Zhang et al., 2025; Bai et al., 2025). Recently, there has been a growing interest in unifying visual understanding and visual generation within a single framework (Team, 2024; Zhou et al., 2024). One line of work tokenizes images into discrete tokens akin to text, enabling large language models (LLMs) to both interpret and generate visual content seamlessly (Team, 2024; Wang et al., 2024). Another approach integrates diffusion processes into LLMs for image generation while relying on additional visual encoders, such as CLIP (Radford et al., 2021; Zhai et al., 2023), for image understanding (Chen et al., 2025; Tong et al., 2024). In the robotics domain, the Unified Video Action Model (Li et al., 2025) proposes a unified architecture that generates images and actions through distinct diffusion heads. In contrast, our WorldVLA explores an alternative direction by employing a discrete autoregressive architecture to build a unified model capable of handling both perception and action generation.

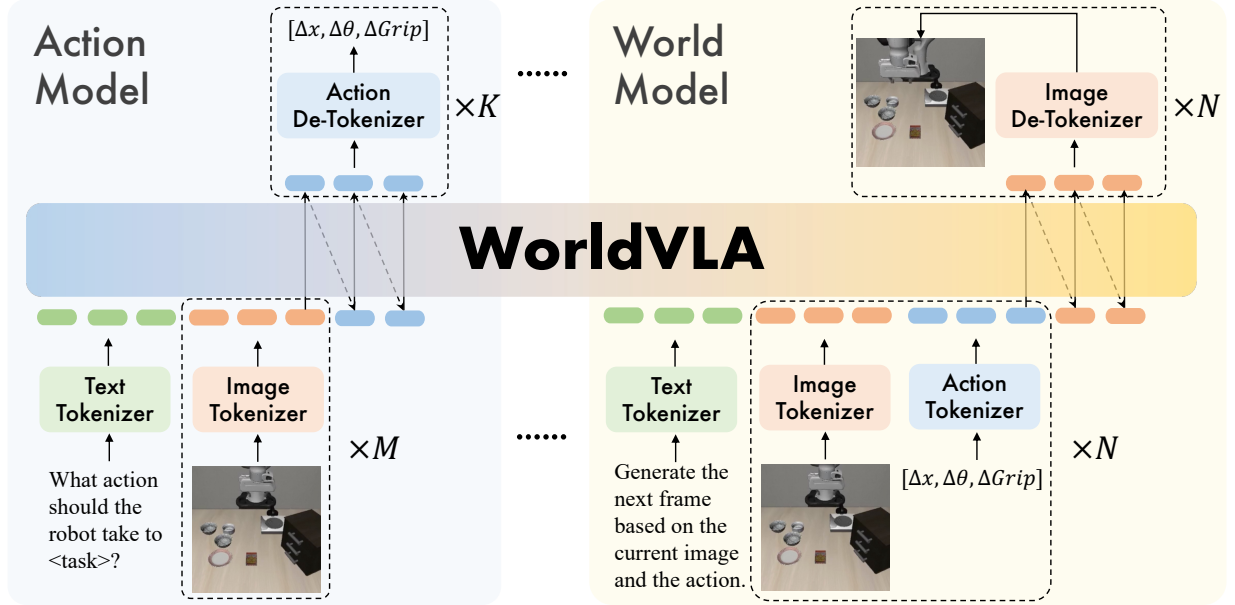


Figure 2 Overview of WorldVLA. WorldVLA integrates two distinct but complementary functional components: an action model and a world model. The action model is responsible for generating actions conditioned on both textual and visual data. The world model functions to predict the subsequent environmental state (e.g., the next visual frame) by leveraging textual information, current image, and current action.

3 Methods

3.1 Problem Formulation

In this work, we address the challenge of learning a unified model capable of simultaneously performing action prediction and world state forecasting. Specifically, we define two primary components: an action model (or policy model) π_θ and a world model f_ϕ . The action model π_θ is responsible for generating an action a_t conditioned on a history of image observations $\{o_{t-h}, o_{t-h+1}, \dots, o_t\}$ and a language instruction l , which can be formally expressed as:

$$a_t = \pi_\theta(a_t \mid o_{t-h:t}, l). \quad (1)$$

Meanwhile, the world model f_ϕ predicts the next frame o_t based on the historical sequence of observations $\{o_{t-h}, o_{t-h+1}, \dots, o_{t-1}\}$ and the corresponding sequence of actions $\{a_{t-h}, a_{t-h+1}, \dots, a_{t-1}\}$. This relationship is formulated as:

$$o_t = f_\phi(o_t \mid o_{t-h:t-1}, a_{t-h:t-1}). \quad (2)$$

Our objective is to develop an integrated action-world model M_ψ that unifies these two functionalities. The model M_ψ should be capable of both predicting actions as a policy model and forecasting future states as a world model. Formally, the unified model M_ψ is defined as:

$$M_\psi : \begin{cases} a_t = M_\psi^{\text{policy}}(a_t \mid o_{t-h:t}, l), \\ o_t = M_\psi^{\text{world}}(o_t \mid o_{t-h:t-1}, a_{t-h:t-1}), \end{cases} \quad (3)$$

where M_ψ^{policy} represents the action generation component and M_ψ^{world} denotes the world state prediction component. By learning such a unified model, we aim to achieve a compact and efficient framework that leverages shared representations for both decision-making and environment modeling.

3.2 Architecture

The overall architecture of autoregressive action world model is shown in Fig. 2. We initialize the model from Chameleon (Team, 2024) since it is a unified model for image understanding and generation. Three tokenizers

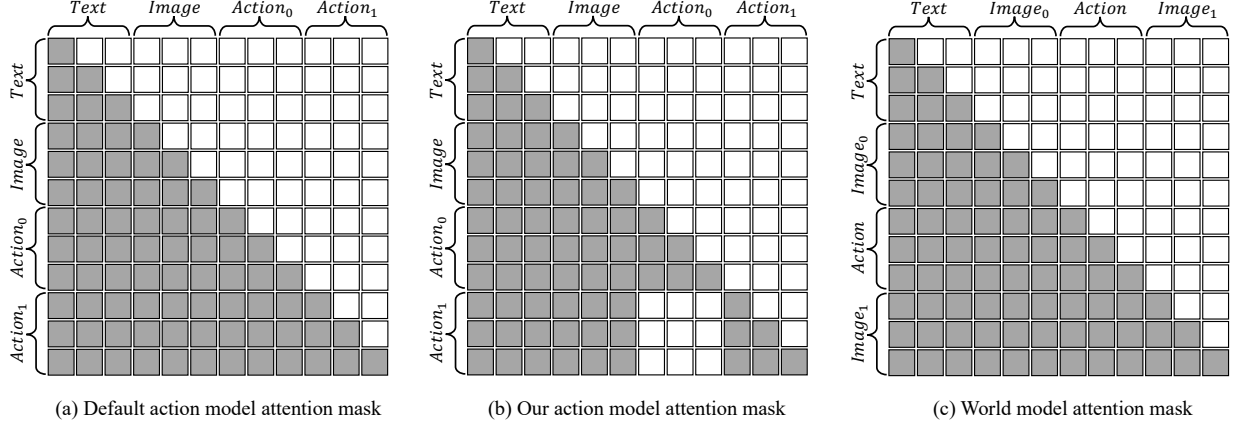


Figure 3 Attention mask mechanism of (a) default action model, (b) our proposed action model, and (c) world model.

are involved, including an image tokenizer, a text tokenizer, and an action tokenizer. The image tokenizer is a VQ-GAN model (Esser et al., 2021) with additional perceptual losses to specific image regions, *e.g.*, faces and salient objects (Gafni et al., 2022). The compression ratio of the image tokenizer is 16 and the codebook size is 8192. The image tokenizer generates 256 tokens for 256×256 images and 1024 tokens for 512×512 images. The action tokenizer discretizes each dimension of continuous robot actions into one of 256 bins, with bin widths determined by the range of the training data (Kim et al., 2024; Brohan et al., 2023). The actions are represented as 7 tokens, including 3 relative positions, 3 relative angles and 1 absolute gripper states. The text tokenizer is a trained BPE tokenizer (Sennrich et al., 2015) with a vocabulary size of 65,536, which includes 8192 image tokens and 256 action tokens. All texts, actions, and images are discretized into tokens and are trained under the autoregressive manner.

3.3 Training Strategy

We mix the action model data and world model data to train our WorldVLA. There are three primary reasons for incorporating world model data to enhance action generation. First, the world model acquires an understanding of the environmental physics by learning to predict future observations based on the current state and applied actions. This learned representation of environmental physics is helpful for manipulation tasks. Second, the world model enables the system to simulate and evaluate potential outcomes of candidate actions, thereby facilitating the avoidance of actions that may lead to unfavorable states. Third, the world model requires a precise interpretation of the action inputs, which in turn supports the action model in producing more effective and contextually appropriate actions. On the other hand, action model enhances visual understanding and in turn supports the visual generation capability of the world model.

Action Model Data. Action model is to generate the action given the text instruction and image observations. The text inputs are *"What action should the robot take to + task instruction + ?"*. The overall token sequence is:

$$[\text{BOS}] \underbrace{\{\text{text}\} [\text{BOI}] \{\text{image}\} \dots \{\text{image}\} [\text{EOI}]}_{\times M} [\text{EOS}] \underbrace{[\text{BOA}] \{\text{action}\} \dots \{\text{action}\} [\text{EOA}]}_{\times K} [\text{EOS}],$$

$\mathcal{L}_{\text{action}}$

where $\{\text{text}\}$, $\{\text{image}\}$, and $\{\text{action}\}$ refer to the discret text, image, and action tokens. $[\text{BOS}]$, $[\text{EOS}]$, $[\text{BOI}]$, $[\text{EOI}]$, $[\text{BOA}]$, $[\text{EOA}]$ refer to the beginning of sentence, end of sentence, beginning of image, end of image tokens, beginning of action, and end of action tokens. The input contains M images and the output contains K actions. We only calculate the loss of action tokens $\mathcal{L}_{\text{action}}$.

World Model Data. World model is to generate the next image frame given the current image observation and action. It does not need the task instruction since the action itself could totally determine the next state. The text inputs are *"Generate the next frame based on the current image and the action."*. The overall token sequence is:

Table 2 Evaluation results on LIBERO benchmark. Pretraining means the model is pretrained on the large-scale robot manipulation data.

Continuous Action Model	Pretraining	Spatial	Object	Goal	Long	Average
Diffusion Policy (Chi et al., 2023)	✗	78.3	92.5	68.3	50.5	72.4
Octo (Team et al., 2024)	✓	78.9	85.7	84.6	51.1	75.1
DiT Policy (Hou et al., 2024)	✓	84.2	96.3	85.4	63.8	82.4
Seer (Tian et al., 2024)	✗	–	–	–	78.7	–
Seer (Tian et al., 2024)	✓	–	–	–	87.7	–
OpenVLA-OFT (Kim et al., 2025)	✓	96.9	98.1	95.5	91.1	95.4
UVA (Li et al., 2025)	✗	–	–	–	93.0	–
Discrete Action Model						
OpenVLA (Kim et al., 2024)	✓	84.7	88.4	79.2	53.7	76.5
WorldVLA (256 * 256)	✗	85.6	89.0	82.6	59.0	79.1
WorldVLA (512 * 512)	✗	87.6	96.2	83.4	60.0	81.8

$$\underbrace{[\text{BOS}]\{\text{text}\} \underbrace{[\text{BOI}]\{\text{image}\}[\text{EOI}][\text{BOA}]\{\text{action}\}[\text{EOA}][\text{EOS}]}_{\times N} \overbrace{[\text{BOI}]\{\text{image}\}[\text{EOI}][\text{EOS}]}^{\mathcal{L}_{\text{world}}}.$$

The next frame prediction conditioned on the action repeats N times, and we only calculate the loss of generated image tokens $\mathcal{L}_{\text{world}}$.

Attention Mask. The standard attention mechanism in autoregressive models typically employs a causal attention mask, which restricts the current token’s access to information exclusively from preceding tokens, excluding any subsequent ones, as illustrated in Fig. 3 (a). Nevertheless, this conventional configuration proves inadequate for generating action chunks, *i.e.*, multiple consecutive actions. While the foundational MLLM demonstrates robust generalization capabilities across image and text domains due to the large-scale pretraining on diverse datasets, its capacity to generalize effectively in the action domain is comparatively limited. Consequently, errors originating from earlier actions propagate to subsequent actions under the default attention mask, resulting in performance degradation. To address this limitation, we introduce an alternative attention mask tailored for action generation, depicted in Fig. 3 (b). This modified mask ensures that current actions rely solely on textual and visual inputs, while prohibiting access to prior actions. Such a design enables the autoregressive framework to generate multiple actions in parallel, aligning with methodologies presented in (Kim et al., 2025; Black et al., 2024). The world model part adheres to the conventional causal attention mask, as shown in Fig. 3 (c).

Training Objective. We mix the action model data and world model data so that the autoregressive action world model could behave as both action model and world model. The loss function is:

$$\mathcal{L} = \mathcal{L}_{\text{action}} + \alpha \mathcal{L}_{\text{world}}, \quad (4)$$

where $\mathcal{L}_{\text{action}}$ and $\mathcal{L}_{\text{world}}$ refer to the cross-entropy loss of the action model data and world model data. Since the image tokens (256 tokens for 256×256 images and 1024 tokens for 512×512 images) are much more than the action tokens (7 tokens), we use α to balance the loss contribution.

4 Experiments

4.1 Evaluation Benchmark

Benchmark. We use LIBERO benchmark (Liu et al., 2023a) in our experiments. LIBERO benchmark contains LIBERO-Spatial, LIBERO-Object, LIBERO-Goal, LIBERO-Long and LIBERO-90. LIBERO-Spatial focuses on spatial relationships by requiring the robot to place a bowl based on its location. LIBERO-Object emphasizes object recognition by having the robot pick and place unique objects. LIBERO-Goal tests

Table 3 Ablation study of action model.

Index	Action Model	World Model	Action Chunking	Our Action Model Attention Mask	Goal SR (%)	Object SR (%)	Spatial SR (%)	Long SR (%)	Average SR (%)
1	✓	✗	✗	✗	67.3	82.9	77.8	23.0	62.8
2	✓	✓	✗	✗	73.1	88.0	80.2	27.3	67.2
3	✓	✗	✓	✗	79.6	82.9	36.7	16.9	54.0
4	✓	✗	✓	✓	84.4	90.9	81.8	49.3	76.6
5	✓	✓	✓	✓	85.1	90.9	84.0	52.4	78.1

procedural learning through varying task goals with fixed objects. LIBERO-Long includes 10 long-horizon tasks. LIBERO-90 provides 90 short-horizon tasks for pretraining.

Datasets. We first filter out the unsuccessful recorded trajectories and no-operation actions like OpenVLA (Kim et al., 2024). Considering the world model evaluation needs ground truth-paired video and action data, we split the 90% of the trajectories as the training set and the 10% remaining trajectories as the validation set. The training set is used for model training by default, with the exception of Table 2, where all available data are utilized during training to ensure a fair comparison.

Baselines. There are two kinds of action models including the continuous action model and discrete action model. Continuous action model generates multiple actions in parallel and uses l1 regression loss for training. Diffusion-based action model like Diffusion Policy (Chi et al., 2023), Octo (Team et al., 2024), DiT Policy (Hou et al., 2024), and UVA (Li et al., 2025) use diffusion process to generate the actions. Seer (Tian et al., 2024) and OpenVLA-OFT (Kim et al., 2025) use an action head to directly output multiple actions in one time. Discrete action models like OpenVLA (Kim et al., 2024) considers the action as tokens just like texts, and the actions are generated in an autoregressive manner. Discrete models inherently exhibit inferior performance, as the tokenization process of actions may lead to information loss.

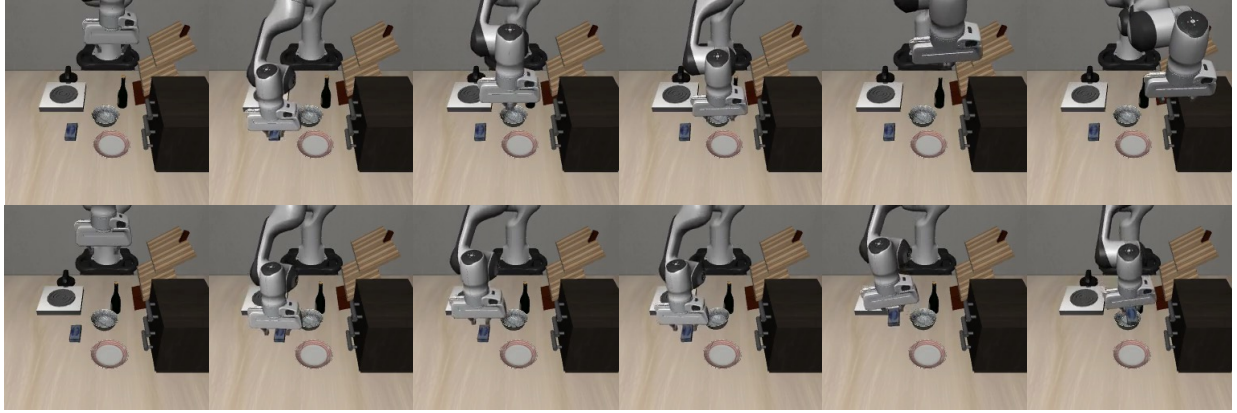
Training Setting. The action model utilizes a default input image count of $M = 2$. The action chunk size is set to $K = 10$ for the LIBERO Long task and $K = 5$ for the remaining three LIBERO tasks under default configuration. To minimize computational expenditure, the world model operates with a single round $N = 1$. The parameter α is fixed at 0.04 in the experimental setup.

Metrics. For action model evaluation, each task is evaluated for 50 rollouts under different initial states and we record the success rates (SR). For world model evaluation, we use the validation set and record the FVD, PSNR, SSIM, and LPIPS values.

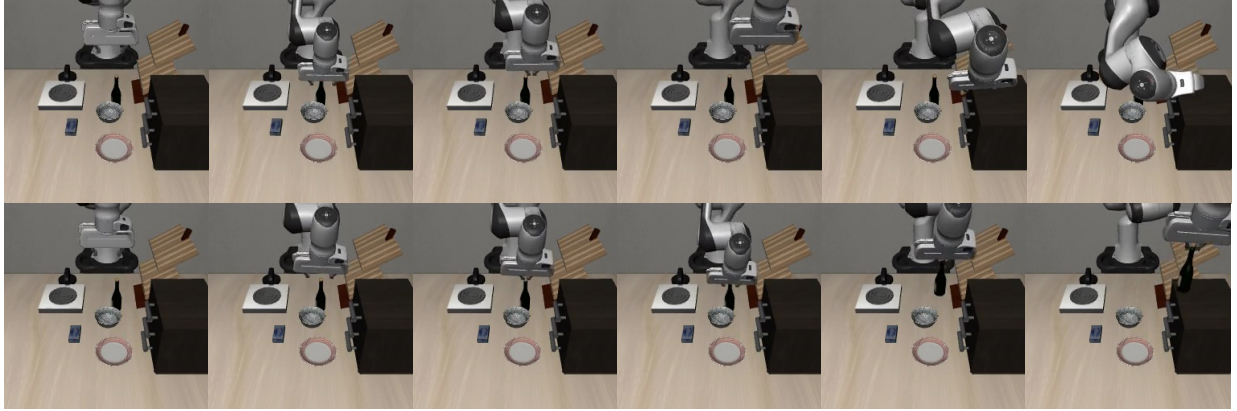
4.2 Evaluation Results and Discussion

Benchmark Results. Table 2 indicates that the proposed WorldVLA model exhibits superior performance compared to the discrete OpenVLA model, even in the absence of pretraining. This outcome shows the effectiveness of the WorldVLA’s design. Furthermore, a positive correlation is observed between image resolution and model performance. Specifically, the $512 * 512$ pixel resolution yielded enhanced results compared to the $256 * 256$ pixel resolution. This phenomenon is primarily attributable to the pretraining regimen of the Chameleon backbone (Team, 2024), whose image tokenization module and the large language model components are inherently optimized at $512 * 512$ resolution. Additionally, higher resolution naturally provides a greater level of detailed visual information, which is particularly crucial for robotic grasping tasks since it demands high operational precision.

World Model Helps Action Model. Quantitative results in Table 3, including row 2 vs. row 1, or row 5 vs. row 4, demonstrate that the integration of a world model significantly enhances the performance of the action model. The world model’s fundamental function involves predicting the subsequent state of the environment, conditioned on the current state and a given action. This generative process inherently promotes the acquisition of an understanding of the system’s underlying physical dynamics, which is a critical prerequisite for successful



(a) Task: put the cream cheese in the bowl. Top: action model. Bottom: our action world model.



(b) Task: put the wine bottle on top of the cabinet. Top: action model. Bottom: our action world model.

Figure 4 Visualization of action model. Top: action model. Bottom: our action world model.

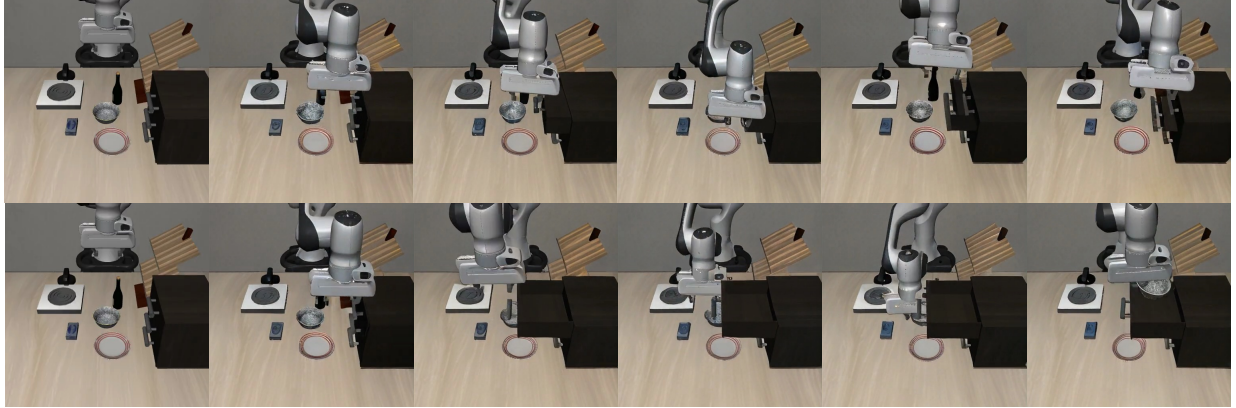
execution in dexterous manipulation tasks such as grasping. Furthermore, the world model endows the system with the capacity for prospective simulation, enabling it to anticipate the consequences of potential actions. This predictive foresight facilitates more informed decision-making, thereby optimizing action selection to maximize the probability of task success. Fig. 4 shows that the action model directly move to the destination without successfully grasping the cheese or bottle. In contrast, our action world model repeatedly attempts to grasp the objects until successful manipulation is achieved before proceeding to the target location.

Action Model Helps World Model. Table 4 demonstrates that the action world model outperforms the pure world model in terms of generation quality, particularly when producing longer video sequences. The action model derives actions based on the input images. On one hand, this contributes to more accurate visual interpretation; on the other, the process of generating actions enhances the understanding of the underlying behavioral patterns. Both aspects support the overall performance of the world model, which relies on robust comprehension of both visual and action-related information to predict future states effectively. As illustrated in Fig. 5, the pure world model fails in several scenarios: it is unable to open the drawer (a), causes the bowl to disappear after moving the disk (b), and fails to lift the bowl onto the stove (c). In contrast, the action world model produces coherent and physically plausible subsequent states in these cases.

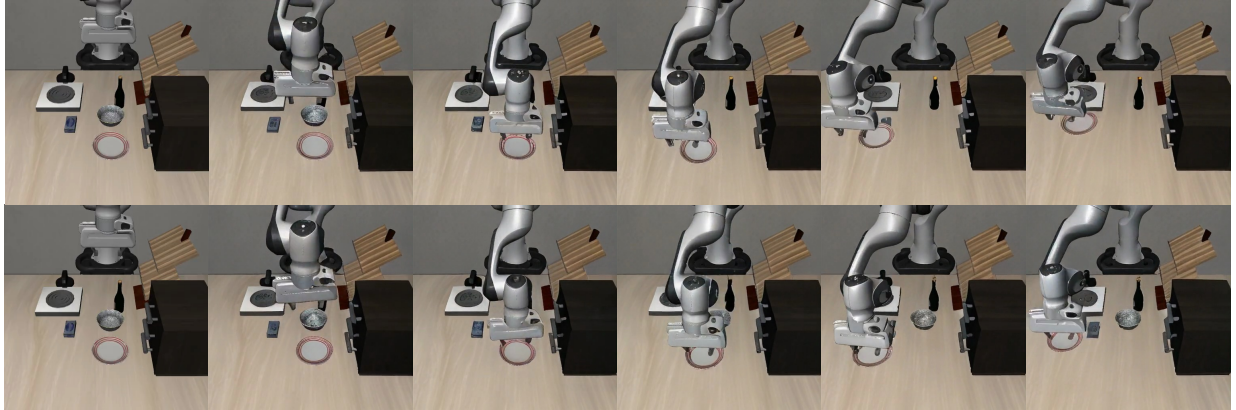
Action Chunking Generation with Proposed Attention Mask. Simultaneous generation of multiple actions is essential for achieving effective and efficient grasping. However, we observe that a naive autoregressive approach—where actions are generated sequentially—can degrade model performance, as evidenced by the results in row 3 of Table 3 and Fig. 6. The grasping success rate gradually decreases with longer action chunks. This degradation arises because later actions become overly dependent on preceding ones since they share the same space, rather than being grounded in visual input which is a distinct modality. The generalization of the action is not that strong as this modality was not involved during pretraining the MLLM. Consequently,

Table 4 Ablation study of world model.

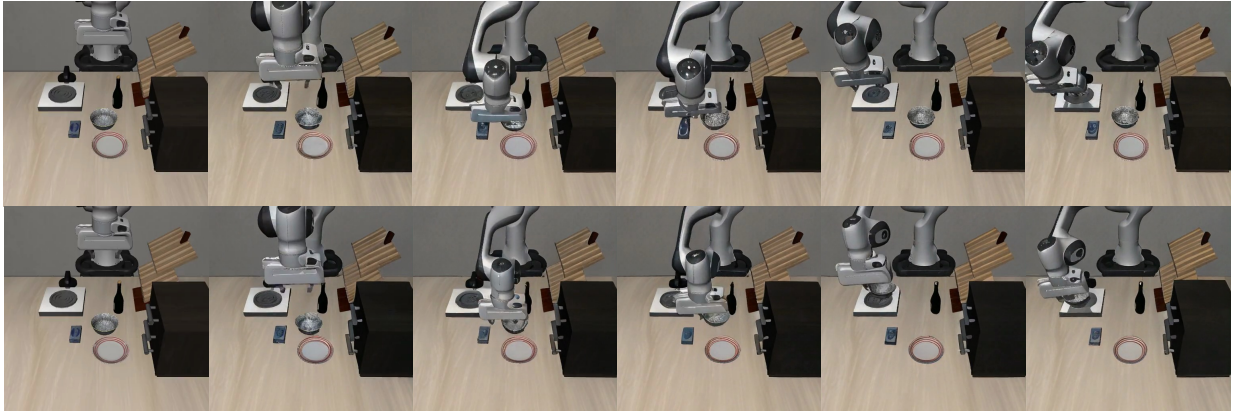
	10 frames				50 frames			
	FVD↓	PSNR↑	SSIM↑	LPIPS↓	FVD↓	PSNR↑	SSIM↑	LPIPS↓
World Model	250.0	29.62	90.73	11.97	718.6	23.98	83.41	15.60
Action World Model	255.1	29.77	90.40	11.94	674.1	24.30	83.55	15.44



(a) Task: open the top drawer and put the bowl inside. Top: world model. Bottom: our action world model.



(b) Task: push the plate to the front of the stove. Top: world model. Bottom: our action world model.



(c) Task: put the bowl on the stove. Top: world model. Bottom: our action world model.

Figure 5 Visualization of world model. Top: world model. Bottom: our action world model.

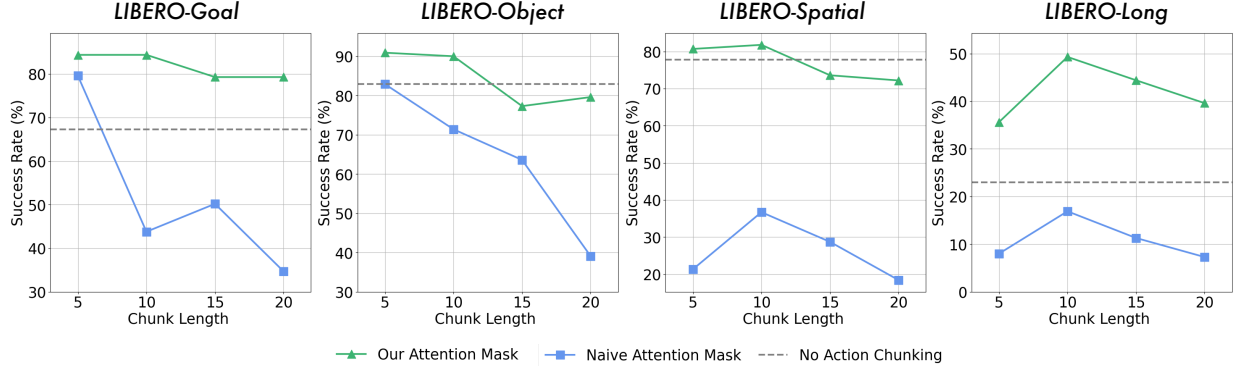


Figure 6 Ablation study of action chunking length.

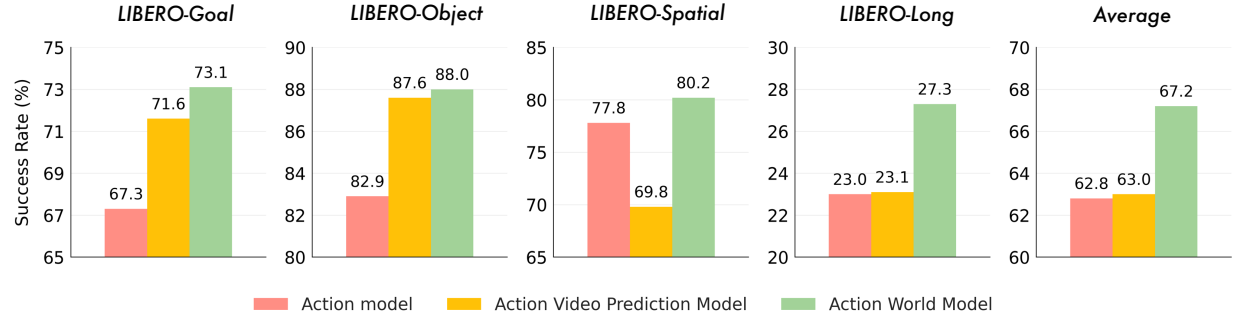


Figure 7 Comparison between action world model and action video prediction model.

errors tend to accumulate as the sequence of generated actions increases. The proposed attention masking mechanism ensures that each action is generated independently and solely determined by the visual input, thereby mitigating the issue of error propagation within the action sequence. As illustrated in Fig. 6, the model incorporating the proposed attention mask demonstrates superior performance compared to the naive attention mask, particularly under conditions of longer chunk lengths. This highlights the efficacy of the introduced masking approach. If the length of the action chunk is excessively prolonged, the robot’s ability to timely adapt its policy becomes constrained, leading to a decline in overall performance, as demonstrated in Fig. 6.

World Model versus Video Prediction Model. Video prediction model is to generate the next frames based on the current frame and the task instruction. Video prediction has been used for pretraining the action model in prior research, such as GR-1 (Wu et al., 2023) and GR-2 (Cheang et al., 2024). Both video prediction model and world model belong to visual generation models, so we conduct a comparison to assess which framework provides greater utility for the action model. The text inputs of video prediction model are *"Generate the future image based on the task and current image. + task instruction"*. The overall token sequence is:

$$[\text{BOS}] \{ \text{text} \} \underbrace{[\text{BOI}] \{ \text{image} \} \dots \{ \text{image} \} [\text{EOI}]}_{\times M} [\text{EOS}] \overbrace{[\text{BOI}] \{ \text{image} \} [\text{EOI}]}^{\mathcal{L}_{\text{video}}} [\text{EOS}].$$

$\times K$

The difference between video prediction model and world model is that the world model is conditioned on the action while video prediction model is not. As illustrated in Fig. 7, the integration of a world model enhances the performance of the action model across all evaluated tasks. The video prediction model, however, demonstrates beneficial effects for two tasks while negatively impacting performance on one task. This discrepancy may arise from the inherent ambiguity in video prediction when action inputs are absent, as the subsequent frame cannot be uniquely determined from the initial frame alone. Consequently, multiple plausible future frames or ground truth sequences may correspond to a single starting frame, potentially introducing noise or inconsistency during training. Furthermore, the incorporation of a world model necessitates an understanding of actions which could contribute to more effective action generation.

Table 5 Ablation study of historical image input length.

	1 frame		2 frames		4 frames	
	SR (%)↑	FPS↑	SR (%)↑	FPS↑	SR (%)↑	FPS↑
w/o Action Chunking	58.4	2.27	67.3	1.77	78.7	1.22
w/ Action Chunking	74.0	3.67	84.4	3.13	84.7	2.78

Table 6 Ablation study of world model pretraining.

	Goal SR (%)	Object SR (%)	Spatial SR (%)	Long SR (%)	Average SR (%)
w/o World Model Pretrain	67.3	82.9	77.8	23.0	62.8
w/ World Model Pretrain	73.1	84.0	79.8	30.2	66.8

Historical Image Input. Unified models for understanding and generation, such as Chameleon (Team, 2024), employ the discrete image tokenizer VQGAN (Esser et al., 2021) for image interpretation. However, their capacity for semantic comprehension is comparatively limited when contrasted with vision-based perceptual models like CLIP (Radford et al., 2021). As demonstrated in Table 5, the use of a single-frame input results in suboptimal performance. To enhance the model’s access to visual context, we incorporate multiple historical image frames, which leads to a progressive improvement in performance. Furthermore, the results indicate that the performance is saturated with two frames when generating action chunks. Consequently, we adopt a two-frame input configuration as the default in our experiments, optimizing the trade-off between task success rate and computational efficiency.

Pretrain Action Model using World Model. Our WorldVLA framework integrates both action model data and world model data during training. We further investigate the possibility of utilizing the world model as a source of pretraining weights for the action model. This form of pretraining necessitates that the model develop an understanding of visual inputs, actions, and the underlying physical dynamics governing state transitions. As presented in Table 6, employing the world model for pretraining leads to notable improvements in grasping performance. These findings highlight the potential of leveraging world model pretraining in robotic applications, particularly in enhancing task-specific performance through prior exposure to general world knowledge.

5 Conclusion and Future Work

This study introduces WorldVLA, a novel autoregressive framework that unifies action and visual understanding with generation capabilities. We demonstrate that the integration of world modeling and action modeling within this architecture can lead to mutual enhancement in performance. An attention mask mechanism has been proposed to enable autoregressive generation of action sequences. Scaling of both data and model size emerges as a promising avenue for further development of the WorldVLA framework. Additionally, the current image tokenizer, which relies on discrete representations, exhibits limitations in perceptual expressiveness; hence, the design of a unified tokenizer capable of both understanding and generating high-quality visual content is an important direction for improvement. The incorporation of an auxiliary action head presents another potential strategy to enhance grasping performance. We anticipate that this work will contribute to and inspire future research in robotics, particularly in the domains of world modeling and unified models for action and image understanding and generation.

References

Niket Agarwal, Arslan Ali, Maciej Bala, Yogesh Balaji, Erik Barker, Tiffany Cai, Prithvijit Chattopadhyay, Yongxin Chen, Yin Cui, Yifan Ding, et al. Cosmos world foundation model platform for physical ai. *arXiv preprint arXiv:2501.03575*, 2025.

- Anurag Ajay, Seungwook Han, Yilun Du, Shuang Li, Abhi Gupta, Tommi Jaakkola, Josh Tenenbaum, Leslie Kaelbling, Akash Srivastava, and Pulkit Agrawal. Compositional foundation models for hierarchical planning. *Advances in Neural Information Processing Systems*, 36:22304–22325, 2023.
- Shuai Bai, Keqin Chen, Xuejing Liu, Jialin Wang, Wenbin Ge, Sibao Song, Kai Dang, Peng Wang, Shijie Wang, Jun Tang, et al. Qwen2. 5-vl technical report. *arXiv preprint arXiv:2502.13923*, 2025.
- Amir Bar, Gaoyue Zhou, Danny Tran, Trevor Darrell, and Yann LeCun. Navigation world models, 2024. <https://arxiv.org/abs/2412.03572>.
- Suneel Belkhale and Dorsa Sadigh. Minivla: A better vla with a smaller footprint, 2024. <https://github.com/Stanford-ILIAD/openvla-mini>.
- Kevin Black, Noah Brown, Danny Driess, Adnan Esmail, Michael Equi, Chelsea Finn, Niccolo Fusai, Lachy Groom, Karol Hausman, Brian Ichter, et al. A vision-language-action flow model for general robot control. *arXiv preprint arXiv:2410.24164*, 2024.
- Andreas Blattmann, Tim Dockhorn, Sumith Kulal, Daniel Mendelevitch, Maciej Kilian, Dominik Lorenz, Yam Levi, Zion English, Vikram Voleti, Adam Letts, et al. Stable video diffusion: Scaling latent video diffusion models to large datasets. *arXiv preprint arXiv:2311.15127*, 2023.
- Anthony Brohan, Noah Brown, Justice Carbajal, Yevgen Chebotar, Joseph Dabis, Chelsea Finn, Keerthana Gopalakrishnan, Karol Hausman, Alex Herzog, Jasmine Hsu, et al. Rt-1: Robotics transformer for real-world control at scale. *arXiv preprint arXiv:2212.06817*, 2022.
- Anthony Brohan, Noah Brown, Justice Carbajal, Yevgen Chebotar, Xi Chen, Krzysztof Choromanski, Tianli Ding, Danny Driess, Avinava Dubey, Chelsea Finn, et al. Rt-2: Vision-language-action models transfer web knowledge to robotic control. *arXiv preprint arXiv:2307.15818*, 2023.
- Qingwen Bu, Jia Zeng, Li Chen, Yanchao Yang, Guyue Zhou, Junchi Yan, Ping Luo, Heming Cui, Yi Ma, and Hongyang Li. Closed-loop visuomotor control with generative expectation for robotic manipulation. *arXiv preprint arXiv:2409.09016*, 2024.
- Jun Cen, Chenfei Wu, Xiao Liu, Shengming Yin, Yixuan Pei, Jinglong Yang, Qifeng Chen, Nan Duan, and Jianguo Zhang. Using left and right brains together: Towards vision and language planning. *arXiv preprint arXiv:2402.10534*, 2024.
- Chi-Lam Cheang, Guangzeng Chen, Ya Jing, Tao Kong, Hang Li, Yifeng Li, Yuxiao Liu, Hongtao Wu, Jiafeng Xu, Yichu Yang, et al. Gr-2: A generative video-language-action model with web-scale knowledge for robot manipulation. *arXiv preprint arXiv:2410.06158*, 2024.
- Xiaokang Chen, Zhiyu Wu, Xingchao Liu, Zizheng Pan, Wen Liu, Zhenda Xie, Xingkai Yu, and Chong Ruan. Janus-pro: Unified multimodal understanding and generation with data and model scaling. *arXiv preprint arXiv:2501.17811*, 2025.
- Cheng Chi, Zhenjia Xu, Siyuan Feng, Eric Cousineau, Yilun Du, Benjamin Burchfiel, Russ Tedrake, and Shuran Song. Diffusion policy: Visuomotor policy learning via action diffusion. *The International Journal of Robotics Research*, page 02783649241273668, 2023.
- Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, et al. An image is worth 16x16 words: Transformers for image recognition at scale. *arXiv preprint arXiv:2010.11929*, 2020.
- Yilun Du, Sherry Yang, Bo Dai, Hanjun Dai, Ofir Nachum, Josh Tenenbaum, Dale Schuurmans, and Pieter Abbeel. Learning universal policies via text-guided video generation. *Advances in neural information processing systems*, 36: 9156–9172, 2023.
- Patrick Esser, Robin Rombach, and Bjorn Ommer. Taming transformers for high-resolution image synthesis. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 12873–12883, 2021.
- Oran Gafni, Adam Polyak, Oron Ashual, Shelly Sheynin, Devi Parikh, and Yaniv Taigman. Make-a-scene: Scene-based text-to-image generation with human priors. In *European Conference on Computer Vision*, pages 89–106. Springer, 2022.
- David Ha and Jürgen Schmidhuber. World models. *arXiv preprint arXiv:1803.10122*, 2018.
- Haoran He, Yang Zhang, Liang Lin, Zhongwen Xu, and Ling Pan. Pre-trained video generative models as world simulators. *arXiv preprint arXiv:2502.07825*, 2025.

- Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 770–778, 2016.
- Zhi Hou, Tianyi Zhang, Yuwen Xiong, Hengjun Pu, Chengyang Zhao, Ronglei Tong, Yu Qiao, Jifeng Dai, and Yuntao Chen. Diffusion transformer policy. *arXiv preprint arXiv:2410.15959*, 2024.
- 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. *arXiv preprint arXiv:2311.12871*, 2023.
- Physical Intelligence, Kevin Black, Noah Brown, James Darpinian, Karan Dhabalia, Danny Driess, Adnan Esmail, Michael Equi, Chelsea Finn, Niccolo Fusai, et al. pi0.5: a vision-language-action model with open-world generalization. *arXiv preprint arXiv:2504.16054*, 2025.
- 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 arXiv:2406.09246*, 2024.
- Moo Jin Kim, Chelsea Finn, and Percy Liang. Fine-tuning vision-language-action models: Optimizing speed and success. *arXiv preprint arXiv:2502.19645*, 2025.
- Bo Li, Yuanhan Zhang, Dong Guo, Renrui Zhang, Feng Li, Hao Zhang, Kaichen Zhang, Peiyuan Zhang, Yanwei Li, Ziwei Liu, et al. Llava-onevision: Easy visual task transfer. *arXiv preprint arXiv:2408.03326*, 2024.
- Shuang Li, Yihuai Gao, Dorsa Sadigh, and Shuran Song. Unified video action model. In *arxiv*, 2025.
- Xinghang Li, Minghuan Liu, Hanbo Zhang, Cunjun Yu, Jie Xu, Hongtao Wu, Chilam Cheang, Ya Jing, Weinan Zhang, Huaping Liu, et al. Vision-language foundation models as effective robot imitators. *arXiv preprint arXiv:2311.01378*, 2023.
- Bo Liu, Yifeng Zhu, Chongkai Gao, Yihao Feng, Qiang Liu, Yuke Zhu, and Peter Stone. Libero: Benchmarking knowledge transfer for lifelong robot learning. *Advances in Neural Information Processing Systems*, 36:44776–44791, 2023a.
- Haotian Liu, Chunyuan Li, Qingyang Wu, and Yong Jae Lee. Visual instruction tuning. *Advances in neural information processing systems*, 36:34892–34916, 2023b.
- Karl Pertsch, Kyle Stachowicz, Brian Ichter, Danny Driess, Suraj Nair, Quan Vuong, Oier Mees, Chelsea Finn, and Sergey Levine. Fast: Efficient action tokenization for vision-language-action models. *arXiv preprint arXiv:2501.09747*, 2025.
- Dean A Pomerleau. Alvin: An autonomous land vehicle in a neural network. *Advances in neural information processing systems*, 1, 1988.
- Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, et al. Learning transferable visual models from natural language supervision. In *International conference on machine learning*, pages 8748–8763. PmLR, 2021.
- David E Rumelhart, Geoffrey E Hinton, and Ronald J Williams. Learning representations by back-propagating errors. *nature*, 323(6088):533–536, 1986.
- Rico Sennrich, Barry Haddow, and Alexandra Birch. Neural machine translation of rare words with subword units. *arXiv preprint arXiv:1508.07909*, 2015.
- Chameleon Team. Chameleon: Mixed-modal early-fusion foundation models. *arXiv preprint arXiv:2405.09818*, 2024.
- Octo Model Team, Dibya Ghosh, Homer Walke, Karl Pertsch, Kevin Black, Oier Mees, Sudeep Dasari, Joey Hejna, Tobias Kreiman, Charles Xu, et al. Octo: An open-source generalist robot policy. *arXiv preprint arXiv:2405.12213*, 2024.
- Yang Tian, Sizhe Yang, Jia Zeng, Ping Wang, Dahua Lin, Hao Dong, and Jiangmiao Pang. Predictive inverse dynamics models are scalable learners for robotic manipulation. *arXiv preprint arXiv:2412.15109*, 2024.
- Shengbang Tong, David Fan, Jiachen Zhu, Yunyang Xiong, Xinlei Chen, Koustuv Sinha, Michael Rabbat, Yann LeCun, Saining Xie, and Zhuang Liu. Metamorph: Multimodal understanding and generation via instruction tuning. *arXiv preprint arXiv:2412.14164*, 2024.
- Xinlong Wang, Xiaosong Zhang, Zhengxiong Luo, Quan Sun, Yufeng Cui, Jinsheng Wang, Fan Zhang, Yueze Wang, Zhen Li, Qiyang Yu, et al. Emu3: Next-token prediction is all you need. *arXiv preprint arXiv:2409.18869*, 2024.

- Junjie Wen, Minjie Zhu, Yichen Zhu, Zhibin Tang, Jinming Li, Zhongyi Zhou, Chengmeng Li, Xiaoyu Liu, Yaxin Peng, Chaomin Shen, et al. Diffusion-vla: Scaling robot foundation models via unified diffusion and autoregression. *arXiv preprint arXiv:2412.03293*, 2024.
- Junjie Wen, Yichen Zhu, Jinming Li, Minjie Zhu, Zhibin Tang, Kun Wu, Zhiyuan Xu, Ning Liu, Ran Cheng, Chaomin Shen, et al. Tinyvla: Towards fast, data-efficient vision-language-action models for robotic manipulation. *IEEE Robotics and Automation Letters*, 2025.
- 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. *arXiv preprint arXiv:2312.13139*, 2023.
- Jialong Wu, Shaofeng Yin, Ningya Feng, Xu He, Dong Li, Jianye Hao, and Mingsheng Long. ivideogpt: Interactive videogpts are scalable world models. *Advances in Neural Information Processing Systems*, 37:68082–68119, 2025.
- Lijun Yu, Yong Cheng, Kihyuk Sohn, José Lezama, Han Zhang, Huiwen Chang, Alexander G Hauptmann, Ming-Hsuan Yang, Yuan Hao, Irfan Essa, et al. Magvit: Masked generative video transformer. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 10459–10469, 2023.
- Xiaohua Zhai, Basil Mustafa, Alexander Kolesnikov, and Lucas Beyer. Sigmoid loss for language image pre-training. In *Proceedings of the IEEE/CVF international conference on computer vision*, pages 11975–11986, 2023.
- Boqiang Zhang, Kehan Li, Zesen Cheng, Zhiqiang Hu, Yuqian Yuan, Guanzheng Chen, Sicong Leng, Yuming Jiang, Hang Zhang, Xin Li, et al. Videollama 3: Frontier multimodal foundation models for image and video understanding. *arXiv preprint arXiv:2501.13106*, 2025.
- Tony Z Zhao, Vikash Kumar, Sergey Levine, and Chelsea Finn. Learning fine-grained bimanual manipulation with low-cost hardware. *arXiv preprint arXiv:2304.13705*, 2023.
- 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. *arXiv preprint arXiv:2403.09631*, 2024.
- Chunting Zhou, Lili Yu, Arun Babu, Kushal Tirumala, Michihiro Yasunaga, Leonid Shamis, Jacob Kahn, Xuezhe Ma, Luke Zettlemoyer, and Omer Levy. Transfusion: Predict the next token and diffuse images with one multi-modal model. *arXiv preprint arXiv:2408.11039*, 2024.