
Show-o2: Improved Native Unified Multimodal Models

Jinheng Xie¹ Zhenheng Yang² Mike Zheng Shou^{1*}

¹ Show Lab, National University of Singapore ² ByteDance

Abstract

This paper presents improved native unified multimodal models, *i.e.*, Show-o2, that leverage autoregressive modeling and flow matching. Built upon a 3D causal variational autoencoder space, unified visual representations are constructed through a dual-path of spatial (-temporal) fusion, enabling scalability across image and video modalities while ensuring effective multimodal understanding and generation. Based on a language model, autoregressive modeling and flow matching are natively applied to the language head and flow head, respectively, to facilitate text token prediction and image/video generation. A two-stage training recipe is designed to effectively learn and scale to larger models. The resulting Show-o2 models demonstrate versatility in handling a wide range of multimodal understanding and generation tasks across diverse modalities, including text, images, and videos. Code and models are released at <https://github.com/showlab/Show-o>.

1 Introduction

Large language models (LLMs) [108, 134] have achieved unprecedented performance levels, fueled by extensive web-scale text resources, substantial computational power, and billions of parameters. In the multimodal domain, large multimodal models (LMMs) [7, 29, 56] and visual generative models [37, 89, 126], have also demonstrated exceptional capabilities in tasks such as general-purpose visual question answering and text-to-image/video generation. Given their success, unified multimodal models (UMMs) [103, 121, 129] have been investigated to unify multimodal understanding and generation within a single model or system. In addition to multimodal understanding capability, this line of approaches seeks to simultaneously cultivate multimodal understanding and generation abilities in the model/system through pre-training, fine-tuning, or connecting tailored models.

Here, we provide a comparative analysis of selected UMMs in Table 1, focusing on two perspectives, including i) visual representations for understanding and generation and ii) the type of unified modeling. Generally, there are two approaches to incorporating visual representations for multimodal understanding and generation: i) a unified representation for both understanding and generation, as seen in works like Chameleon [103], Transfusion [148], and Show-o [129]; and ii) decoupled representations, utilizing CLIP [91] for multimodal understanding and variational autoencoder (VAE) for visual generation. To involve both multimodal understanding and generation capabilities, two primary methods have been explored: i) natively applying multimodal understanding and generation objectives within a single model and ii) tuning adapters to assemble tailored models. We refer the first type as *native unified multimodal models*, distinguishing it from the second type that assembles tailored models. These principles, combined with autoregressive or diffusion modeling or both, contribute to the development of unified multimodal models.

Compared to existing UMMs that primarily focus on text and image, our approach explores model designs that provide substantial potential and scalability in natively unifying text, image, and video modalities. An overview of our approach is presented in Fig. 1. Specifically, for visual inputs, we

* Corresponding Author

Table 1: Comparative analysis of selected unified multimodal models based on the type of visual representations and unified modeling for multimodal understanding and generation. In this context, **native und. & gen.** refers to the direct decoding of output sequences into texts, images, and videos, as opposed to serving as conditions for decoding using external pre-trained decoders like Stable Diffusion. * indicates the method adopts two distinct models for multimodal understanding and generation, respectively. Diff. means the diffusion modeling. *Please refer to the complete Table 15 in the appendix.*

Methods	Und. & Gen. Representation			Type of Unified Modeling		
	Unified	Decoupled	Support Video	Native Und. & Gen.	Assembling Tailored Models	Paradigm
Chameleon [103]	✓		✗	✓		AR
Transfusion [148]	✓		✗	✓		AR + Diff.
Show-o [129]	✓		✗	✓		AR + Diff.
VILA-U [124]	✓		✓	✓		AR
Emu3 [115]	✓		✓	✓		AR
LlamaFusion [96]	✓		✗	✓		AR + Diff.
Show-o2 (Ours)	✓		✓	✓		AR + Diff.
Janus-Series [26, 27, 80]		✓	✗	✓		AR (+Diff)
UnidFluid [38]		✓	✗	✓		AR + MAR
Mogao [66]		✓	✗	✓		AR + Diff.
BAGEL [32]		✓	✓	✓		AR + Diff.
NExT-GPT [121]		✓	✓		✓	AR + Diff.
SEED-X [40]		✓	✗		✓	AR + Diff.
ILLUME [112]		✓	✗		✓	AR + Diff.
MetaMorph [107]		✓	✗		✓	AR + Diff.
MetaQueries [84]		✓	✗		✓	AR + Diff.
TokenFlow* [90]	✓		✗		✓	AR

operate within the 3D causal VAE [109] space, which is capable of accommodating both images and videos. Recognizing the distinct feature dependencies between multimodal understanding and generation, we construct unified visual representations that simultaneously capture rich semantic information and low-level features with intrinsic structures and textual details from the visual latents. This is achieved through a dual-path mechanism consisting of semantic layers, a projector, and a spatial (-temporal) fusion process. As the fusion process occurs within the 3D causal VAE space, when it comes to videos, semantic and low-level features are temporally aligned and fused with full-frame video information.

Text embeddings and unified visual representations are structured into a sequence to go through a pre-trained language model and are modeled by a specific language head and flow head, respectively. Specifically, autoregressive modeling with causal attention is performed on the language head when dealing with text token prediction, and flow matching with full attention is applied to the flow head for image/video generation. Since the base language model lacks visual generation capabilities, we propose a two-stage training recipe to effectively learn such an ability while retaining the language knowledge, without requiring a massive text corpus. In the first stage, we mainly focus on pre-training the flow head for visual generation using (interleaved) text, image, and video data. In the second stage, the full model is fine-tuned with high-quality multimodal understanding and generation data.

Extensive experimental results have demonstrated that our model surpasses the existing methods in terms of most metrics across multimodal understanding and visual generation benchmarks. Collectively, the main contributions of this paper can be summarized as:

- We present an improved native unified multimodal model that seamlessly integrates autoregressive modeling and flow matching, enabling a wide range of multimodal understanding and generation across (interleaved) text, images, and videos.
- Based on the 3D causal VAE space, we construct unified visual representations scalable to both multimodal understanding and generation, image and video modalities by combining semantic and low-level features through a dual-path of spatial (-temporal) fusion mechanism.
- We design a two-stage training pipeline that effectively and efficiently learns unified multimodal models, retaining language knowledge and enabling effective scaling up to larger models, without requiring a massive text corpus.
- The proposed model demonstrates state-of-the-art performance on multimodal understanding and visual generation benchmarks, surpassing existing methods across various metrics.

2 Related Work

2.1 Large Multimodal Models

Building upon the advancements of large language models (LLMs) [108, 134], large multimodal models (LMMs) [7, 29, 56, 61, 73] have showcased remarkable capabilities in general-purpose visual question answering. These approaches typically leverage pre-trained vision encoders to project visual features and align them within the embedding space of LLMs. Meanwhile, a growing number of encoder-free LMMs [34, 35, 129] aim to directly align raw visual features within the LLM embedding space. However, these encoder-free methods often fall behind models that utilize image-text-aligned visual features in terms of performance. Beyond model architecture, recent studies [20, 56, 106] have highlighted the critical role of high-quality instructional data in enhancing multimodal capabilities.

2.2 Visual Generative Models

Two prominent paradigms for visual generation, namely diffusion [9, 19, 71, 87, 89, 93, 94, 120, 126, 128, 141] and autoregressive modeling [22, 52, 59, 85, 98], have been extensively studied in image and video generation in recent years. Diffusion-based methods typically employ optimized architectures that integrate pre-trained text encoders with denoising networks. In contrast, autoregressive methods often utilize LLM-based architectures and are trained through next-token prediction. Recently, several studies [38, 62, 74] have explored hybrid approaches that combine diffusion and autoregressive modeling to further advance visual generation capabilities.

2.3 Unified Multimodal Models

Building on the success of large multimodal and visual generative models, pioneering unified multimodal models (UMMs) such as Chameleon [103], Show-o [129], and Transfusion [148] aim to integrate these capabilities into a single model through autoregressive or diffusion modeling or both. Further advancements [28, 50, 78, 97, 115, 124] have focused on optimizing the training pipeline and enhancing the semantics of discrete tokens, leading to improved performance. We refer to these approaches as *native unified multimodal models*, as they inherently combine multimodal understanding and generation objectives within a unified architecture.

An alternative and promising direction [18, 36, 40, 77, 84, 102, 107] for unifying multimodal understanding and generation involves assembling off-the-shelf specialized LMMs and visual generative models by tuning adapters or learnable tokens. Representative works [18, 40, 84, 121] have demonstrated the promising capabilities and intriguing properties of such assembled unified frameworks, highlighting their potential for further exploration.

3 Methodology

In this section, we introduce the overall framework (Section 3.1), which consists of two key components: i) the design of unified visual representations for multimodal understanding and generation, applicable to both images and videos, and ii) the native learning of multimodal understanding and generation capabilities. Subsequently, we present a two-stage training recipe (Section 3.2), which is designed to progressively learn and effectively scale up the unified multimodal model.

3.1 Overall Framework

Overall Architecture. An overview of our proposed unified model is depicted in Fig. 1. Given (interleaved) texts, images, or videos, a text tokenizer with an embedding layer and a 3D causal VAE encoder accordingly process them into continuous text embeddings and visual latent representations. Subsequently, the visual latent representations undergo a dual-path extraction of spatial (-temporal) fusion to create the unified visual representations. These representations are then structured into a sequence, which is fed into a language model equipped with language and flow heads to model the sequence via autoregressive modeling and flow matching accordingly. Finally, a text de-tokenizer in conjunction with a 3D causal VAE decoder is employed to decode the final output. Next, we will delve into the fundamental design principles behind the unified visual representation and flow head.

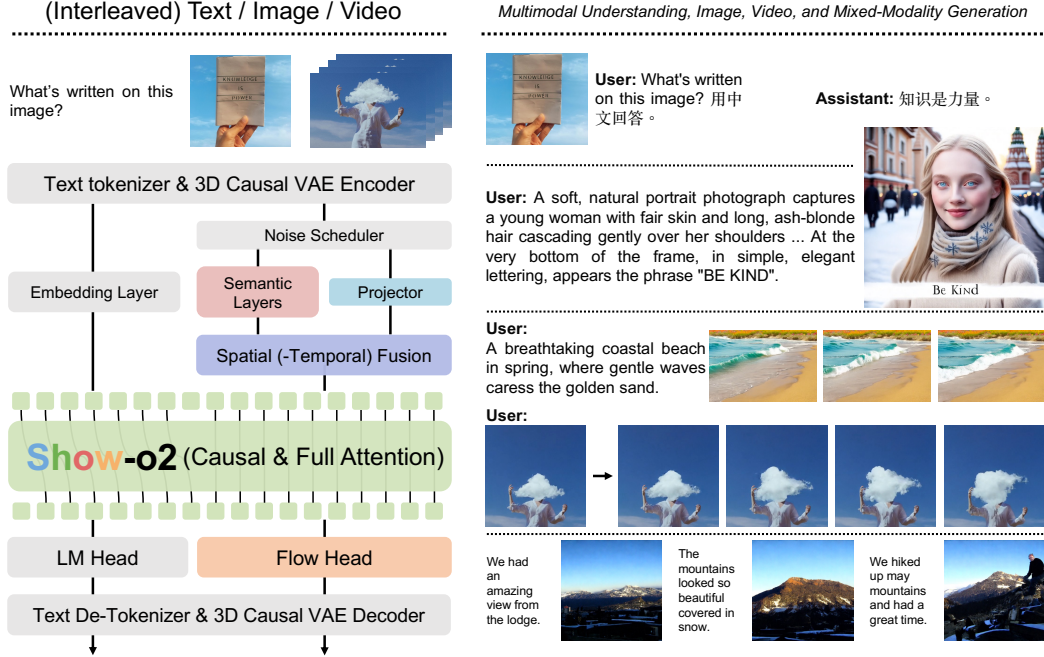


Figure 1: Our approach begins by encoding input texts, images, and videos into continuous embeddings and visual latents. The visual latents are processed through a dual-path extraction and spatial (-temporal) fusion mechanism to construct unified visual representations that are scalable for both multimodal understanding and generation, image and video modalities. These text embeddings and unified visual representations are then structured into a sequence for the base language model, equipped with dedicated heads. Specifically, text tokens are modeled autoregressively by a language head, while image and video latents are handled by a flow head using flow matching. We employ the omni-attention mechanism [129, 148] to enable causal attention along the sequence while maintaining full attention within the unified visual representations. This design empowers our model to effectively tackle tasks such as image/video understanding, generation, and mixed-modality generation.

Unified Visual Representation. To scalably support image and video modalities, we employ a 3D causal VAE encoder to extract image/video latents. As multimodal understanding and generation differ in feature dependency, we propose a dual-path architecture comprising semantic layers $\mathcal{S}(\cdot)$ to extract high-level representations of rich semantic contextual information and a projector $\mathcal{P}(\cdot)$ to retain complete low-level information from the extracted visual latents. Specifically, semantic layers $\mathcal{S}(\cdot)$ share the same vision transformer blocks of SigLIP [139] with a new 2×2 patch embedding layer. Given n visual latents $\mathbf{x}_t = \{x_i\}_{i=1}^n$ at a noise level:

$$\mathbf{x}_t = t \cdot \mathbf{x}_1 + (1 - t) \cdot \mathbf{x}_0, \quad (1)$$

where $\mathbf{x}_0 \sim \mathcal{N}(0, 1)$ and $t \sim [0, 1]$, we load the pre-trained weights of SigLIP and pre-distill $\mathcal{S}(\cdot)$ as follows:

$$\mathcal{L}_{\text{distill}} = -\frac{1}{n} \sum \log \text{sim}(\mathcal{S}(\mathbf{x}_t), \text{SigLIP}(\mathbf{X})), \quad (2)$$

where $\mathbf{X}$ is the input image, $\text{SigLIP}(\cdot)$ extracts the image patch features, and $\text{sim}(\cdot)$ indicates the cosine similarity calculator. In this way, semantic layers $\mathcal{S}(\cdot)$ can mimic extracting semantic features from both clean and noised visual latents $\mathbf{x}_t$. The projector $\mathcal{P}(\cdot)$ is simply composed of a 2D patch embedding layer. The extracted high- and low-level representations are spatially (and temporally when it comes to videos) fused by concatenating through the feature dimension and applying RMSNorm [140] with two MLP layers to get the unified visual representations $\mathbf{u}$:

$$\mathbf{u} = \text{STF}(\mathcal{S}(\mathbf{x}_t), \mathcal{P}(\mathbf{x}_t)), \quad (3)$$

where STF indicates the spatial (-temporal) fusion mechanism. In addition, we prepend a time step t embedding to the unified visual representations for generative modeling. t is set as 1.0 to get time step embedding for the clean image.

We structure the text embeddings and unified visual representations into a sequence following a general interleaved image-text format below:

$$[\text{BOS}] \{\text{Text}\} [\text{BOI} / \text{BOV}] \{\text{Image} / \text{Video}\} [\text{EOI} / \text{EOV}] \{\text{Text}\} \cdots [\text{EOS}].$$

The sequence format above is flexible and can be adapted to various input types. We adopt the omni-attention mechanism [129, 148] to let the sequence modeling be causal but with full attention within the unified visual representations.

Flow Head. Apart from the language head for text token prediction, we employ a flow head to predict the defined velocity $\mathbf{v}_t = \frac{d\mathbf{x}_t}{dt}$ via flow matching [71, 75]. Specifically, the flow head simply consists of several transformer layers with time step modulation via the adaLN-Zero blocks, as seen in DiT [86].

During training, we natively apply next token prediction $\mathcal{L}_{\text{NTP}}$ to the language head and flow matching $\mathcal{L}_{\text{FM}}$ to the flow head for predicting velocity, respectively:

$$\mathcal{L} = \alpha \mathcal{L}_{\text{NTP}} + \mathcal{L}_{\text{FM}}. \quad (4)$$

3.2 Training Recipe

Existing UMMs, such as Show-o [129], Janus-Pro [26], Transfusion [148], Chameleon [103], and Emu3 [115], are typically trained from LLMs, LMMs, or from scratch. These approaches aim to cultivate

Table 2: Trainable components and datasets in the training stages.

Trainable Components		Datasets		
		# Image-Text	# Video-Text	# Interleaved Data
Stage-1	Projector Spatial (-Temporal) Fusion Flow Head	66M	WebVid [8] Pandas [23]	OmniCorpus [60]
Stage-2	Full Model (w/o VAE)	9M HQ Und. 16M HQ Gen.	OpenVid-1M [81] Gen. 1.5M Internal Data Gen. 1.6M Video Und.	VIST [47] CoMM [24]

visual generative modeling capabilities while preserving language modeling proficiency. However, this process often relies on web-scale, high-quality text corpora, which are prohibitively expensive to collect. Consequently, the lack of such resources can lead to a degradation in language knowledge and modeling performance. To address this challenge, we adopt a two-stage training recipe (as shown in Table 2) that effectively retains language knowledge while simultaneously developing visual generation capabilities, without requiring a massive text corpus.

Stage-1. Before the two-stage training, we have pre-distilled the semantic layers $\mathcal{S}(\cdot)$ (implementation details can be found in Section 4). The first stage only involves trainable components of the projector, spatial (-temporal) fusion, and flow head. In this stage, we train these components using autoregressive modeling and flow matching using around 66M image-text pairs and progressively add interleaved data and video-text pairs.

Stage-2. Subsequently, we tune the full model using 9M high-quality multimodal understanding instruction data, 16M high-quality visual generation data filtered from the 66M image-text pairs, and 1.6M video understanding data.

Scaling Up. After the training of the small-sized model with approximately 1.5B LLM parameters, we resume the pre-trained flow head for the larger model with 7B LLM parameters and introduce a lightweight MLP transformation to align the hidden size, allowing it to quickly adapt to the larger model and converge.

4 Experiments

Dataset curation and implementation details are provided in the Appendix A.

4.1 Multimodal Understanding on Images and Videos

Quantitative Results. Table 3 highlights the performance of our models on multimodal understanding benchmarks, evaluated across metrics such as MME [39], GQA [49], SEED-Bench [55], MM-Bench [76], MMU [138], MMStar [21], and AI2D [51]. As shown in the table, both the 1.5B and 7B variants of our model consistently outperform state-of-the-art models across many metrics. For models with similar parameter sizes (1.5B), our model achieves the best scores on MME-p and MMU-val benchmarks while delivering competitive performance on GQA and SEED-Bench metrics. When

Table 3: Evaluation on multimodal understanding benchmarks. # Params. indicates the number of parameters of base LLM. * indicates the method uses two distinct models or sets of parameters for multimodal understanding and generation, respectively. † indicates the Show-o2 models fine-tuned using video understanding data. Und. indicates “understanding”. Results in gray indicate the performance of und. only models or models with total parameters more than 13B.

Types	Models	# Params.	MME ↑ (p)	GQA ↑	SEED ↑ (all)	MMB ↑ (en)	MMMU ↑ (val)	MMStar ↑	AI2D ↑
Und. Only	LLaVA-v1.5 [72]	7B	1510.7	62.0	58.6	64.3	-	-	-
	Qwen-VL-Chat [6]	7B	1487.6	57.5	58.2	60.6	-	-	57.7
	LLaVA-OV [56]	7B	1580.0	-	-	80.8	48.8	57.5	81.4
Unify via Assembling Tailored Models	NExT-GPT [129]	13B	-	-	57.5	58.0	-	-	-
	SEED-X [40]	17B	1457.0	49.1	66.5	70.1	35.6	-	-
	MetaMorph [107]	8B	-	-	71.8	75.2	-	-	-
	TokenFlow-XL* [90]	14B	1551.1	62.5	72.6	76.8	43.2	-	75.9
	ILLUME [112]	7B	1445.3	-	72.9	75.1	38.2	-	71.4
Native Unified	BAGEL [32]	14B	1687.0	-	-	85.0	55.3	-	-
	Show-o [129]	1.3B	1097.2	58.0	51.5	-	27.4	-	-
	JanusFlow [80]	1.5B	1333.1	60.3	70.5	74.9	29.3	-	-
	SynerGen-VL [58]	2.4B	1381.0	-	-	53.7	34.2	-	-
	Janus-Pro [26]	1.5B	1444.0	59.3	68.3	75.5	36.3	-	-
	Show-o2 (Ours)	1.5B	1450.9	60.0	65.6	67.4	37.1	43.4	69.0
	Emu3 [115]	8B	-	60.3	68.2	58.5	31.6	-	70.0
	VILA-U [124]	7B	1401.8	60.8	59.0	-	-	-	-
	MUSE-VL [130]	7B	-	-	69.1	72.1	39.7	49.6	69.8
	Liquid [119]	8B	1448.0	61.1	-	-	-	-	-
	Janus-Pro [26]	7B	1567.1	62.0	72.1	79.2	41.0	-	-
	Mogao [66]	7B	1592.0	60.9	74.6	75.0	44.2	-	-
	Show-o2 (Ours)	7B	1620.5	63.1	69.8	79.3	48.9	56.6	78.6

Table 4: Evaluation on video understanding benchmarks. # Params. denotes the number of parameters in the base LLM, while # Frames represents the maximum number of video frames used during training and inference. Und. stands for understanding. † marks the Show-o2 models that have been fine-tuned on video understanding data. All results are reported in terms of zero-shot accuracy.

Model	# Params.	# Frames	ActNet-QA	MVBench	NExT-QA	PerceptionTest	LongVideoBench	VideoMME
			test	test	mc	val	val	wo/w-sub
Proprietary Und. Only Models								
GPT-4V [82]	-	-	57.0	43.5	-	-	61.3	59.9/63.3
GPT-4o [83]	-	-	-	-	-	-	66.7	71.9/77.2
Gemini-1.5-Flash [104]	-	-	55.3	-	-	-	61.6	70.3/75.0
Gemini-1.5-Pro [104]	-	-	57.5	-	-	-	64.0	75.0/81.3
Open-source Und. Only Models								
VILA [70]	40B	-	58.0	-	67.9	54.0	-	60.1/61.1
PLLaVA [132]	34B	16 / 16	60.9	58.1	-	-	53.2	-
LongVA [144]	7B	-	50.0	-	68.3	-	-	52.6/54.3
IXC-2.5 [143]	7B	64 / 64	52.8	69.1	71.0	34.4	-	55.8/58.8
LLaVA-OV [56]	7B	32 / 32	56.6	56.7	79.4	57.1	56.5	58.2/61.5
VideoLLaMA2 [30]	7B	16 / 16	50.2	54.6	-	51.4	-	47.9/50.3
Unified Multimodal Models								
Show-o2 [†] (Ours)	1.5B	32 / 32	52.7	49.8	72.1	56.1	49.2	48.0/51.6
Show-o2 [†] (Ours)	7B	16 / 32	56.4	55.8	79.0	61.9	55.5	57.4/60.9

compared to larger models with approximately 7B parameters, our models surpass state-of-the-art models such as Janus-Pro and even the significantly larger TokenFlow-XL model (14B parameters) in metrics including MME-p, GQA, MMMU-val, MMStar, and AI2D, while maintaining competitive performance on SEED-Bench and MM-Bench. In addition, we present the video understanding performance of Show-o2† in Table 4.

Qualitative Results. Fig. 2 showcases the multimodal understanding capabilities of our model. As demonstrated, the model excels at answering general-purpose questions about an image. Specifically, it can provide detailed descriptions of an image, count objects, and recognize text within the image. Besides, the model can leverage its world knowledge to offer step-by-step instructions for preparing daily drinks like an avocado milkshake and supports bilingual question-answering, highlighting its versatility and practical utility.

Table 5: Evaluation on the GenEval [41] benchmark. Gen. denotes “generation”. # Params. indicates the number of parameters of base LLM. # Data. indicates the number of image-text pairs used for visual generation during training. * means the method uses two distinct models for multimodal understanding and generation, respectively. Obj.: Object. Attri.: Attribute. Our results are obtained using rewritten prompts. + indicates the additional data required by the pretrained diffusion models.

Type	Method	# Params.	# Data	Single Obj.	Two Obj.	Counting	Colors	Position	Color Attri.	Overall↑
Gen. Only	SD3-Medium [37]	-	-	0.99	0.94	0.72	0.89	0.33	0.60	0.74
Unifying via Assembling Tailored Models	SEED-X [40]	17B	158M+	0.97	0.58	0.26	0.80	0.19	0.14	0.49
	TokenFlow-XL* [90]	14B	60M	0.95	0.60	0.41	0.81	0.16	0.24	0.55
	ILLUME [112]	7B	15M+	0.99	0.86	0.45	0.71	0.39	0.28	0.61
	MetaQuery-XL [84]	7B	28M+	-	-	-	-	-	-	0.80
Native Unified	Show-o [129]	1.3B	2.0B	0.98	0.80	0.66	0.84	0.31	0.50	0.68
	Emu3 [115]	8B	-	-	-	-	-	-	-	0.66
	MUSE-VL [130]	7B	24M	-	-	-	-	-	-	0.57
	Transfusion [148]	7B	3.5B	-	-	-	-	-	-	0.63
	D-DiT [64]	2B	40M	0.97	0.80	0.54	0.76	0.32	0.50	0.65
	Janus-Pro [26]	7B	144M	0.99	0.89	0.59	0.90	0.79	0.66	0.80
	BAGEL [32]	14B	1600M	0.98	0.95	0.84	0.95	0.78	0.77	0.88
	Mogao [66]	7B	-	1.00	0.97	0.83	0.93	0.84	0.80	0.89
	Show-o2 (Ours)	1.5B	66M	0.99	0.86	0.55	0.86	0.46	0.63	0.73
	Show-o2 (Ours)	7B	66M	1.00	0.87	0.58	0.92	0.52	0.62	0.76

Table 6: Evaluation on the DPG-Bench [45] benchmark. Gen. denotes “generation”. # Params. indicates the number of parameters of base LLM. # Data. indicates the number of image-text pairs used for visual generation during training.

Type	Method	# Params.	# Data	Global	Entity	Attribute	Relation	Other	Overall↑
Gen. Only	Hunyuan-DiT [65]	1.5B	-	84.59	80.59	88.01	74.36	86.41	78.87
	Playground v2.5 [57]	-	-	83.06	82.59	81.20	84.08	83.50	75.47
	PixArt-Σ [17]	-	-	86.89	82.89	88.94	86.59	87.68	80.54
	DALL-E 3 [10]	-	-	90.97	89.61	88.39	90.58	89.83	83.50
	SD3-Medium [37]	2B	-	87.90	91.01	88.83	80.70	88.68	84.08
Native Unified	Emu3-DPO [115]	8B	-	-	-	-	-	-	81.60
	Janus-Pro [26]	7B	144M	86.90	88.90	89.40	89.32	89.48	84.19
	Mogao [66]	7B	-	82.37	90.03	88.26	93.18	85.40	84.33
	Show-o2 (Ours)	1.5B	66M	87.53	90.38	91.34	90.30	91.21	85.02
	Show-o2 (Ours)	7B	66M	89.00	91.78	89.96	91.81	91.64	86.14

Table 7: Overall quantitative comparison of different methods on OneIG-Bench. Gen. denotes “generation”. # Params. indicates the number of parameters of base LLM. # Data. indicates the number of image-text pairs used for visual generation during training.

Type	Method	# Params.	# Data	Alignment↑	Text↑	Reasoning↑	Style↑	Diversity↑
Gen. Only	SD3.5-Large [37]	8B	-	0.809	0.629	0.294	0.353	0.225
	Flux.1-dev [54]	12B	-	0.786	0.523	0.253	0.368	0.238
	SANA-1.5 (PAG) [127]	4.8B	-	0.765	0.069	0.217	0.401	0.216
	Lumina-Image 2.0 [89]	2.6B	110M	0.819	0.106	0.270	0.354	0.216
	HiDream-I1-Full [44]	17B	-	0.829	0.707	0.317	0.347	0.186
Unified Models	Show-o-512 [129]	1.3B	2B	0.702	0.002	0.213	0.361	0.241
	Janus-Pro [27]	7B	144M	0.553	0.001	0.139	0.276	0.365
	BLIP3-o [18]	8B	55M	0.711	0.013	0.223	0.361	0.229
	BAGEL [32]	14B	1600M	0.769	0.244	0.173	0.367	0.251
	OmniGen2 [118]	7B	150M	0.804	0.680	0.271	0.377	0.242
	Show-o2 (Ours)	1.5B	66M	0.798	0.002	0.219	0.317	0.186
	Show-o2-1024×1024 (Ours)	1.5B	66M	0.798	0.125	0.274	0.351	0.186
	Show-o2 (Ours)	7B	66M	0.817	0.002	0.226	0.317	0.177

4.2 Visual Generation

Image Generation. We compare our model with the state-of-the-art approaches on GenEval [41], DPG-Bench [45], and OneIG [13] benchmarks in Tables 5, 6, and 7. One can observe that our model surpasses most approaches, including TokenFlow-XL, Show-o, Emu3, and Transfusion, on the GenEval benchmark. Compared to Janus-Pro, which was trained on a significantly larger dataset of 144M image-text pairs, our model achieves promising results with only 66M image-text pairs. On DPG-Bench evaluation, our model has demonstrated the best overall score compared to generation-only models such as SD3-Medium and unified models, including Emu3-DPO and Janus-Pro. On OneIG-Bench, our models also achieve competitive performance. We also show qualitative results in Fig. 2 to illustrate that our model can generate high-quality and realistic images.

Video Generation. We compare our model with the text-to-video and image-to-video generation models in Tables 8 and 9. One can observe that with only 2B parameters, our model outperforms models such as Show-1, Emu3, and VILA-U with more than 6B parameters. Besides, our model has demonstrated competitive performance compared to CogVideoX and Step-Video-T2V. We also provide qualitative results of the text-to-video and image-to-video generation capability of our model in the middle of Fig. 2.

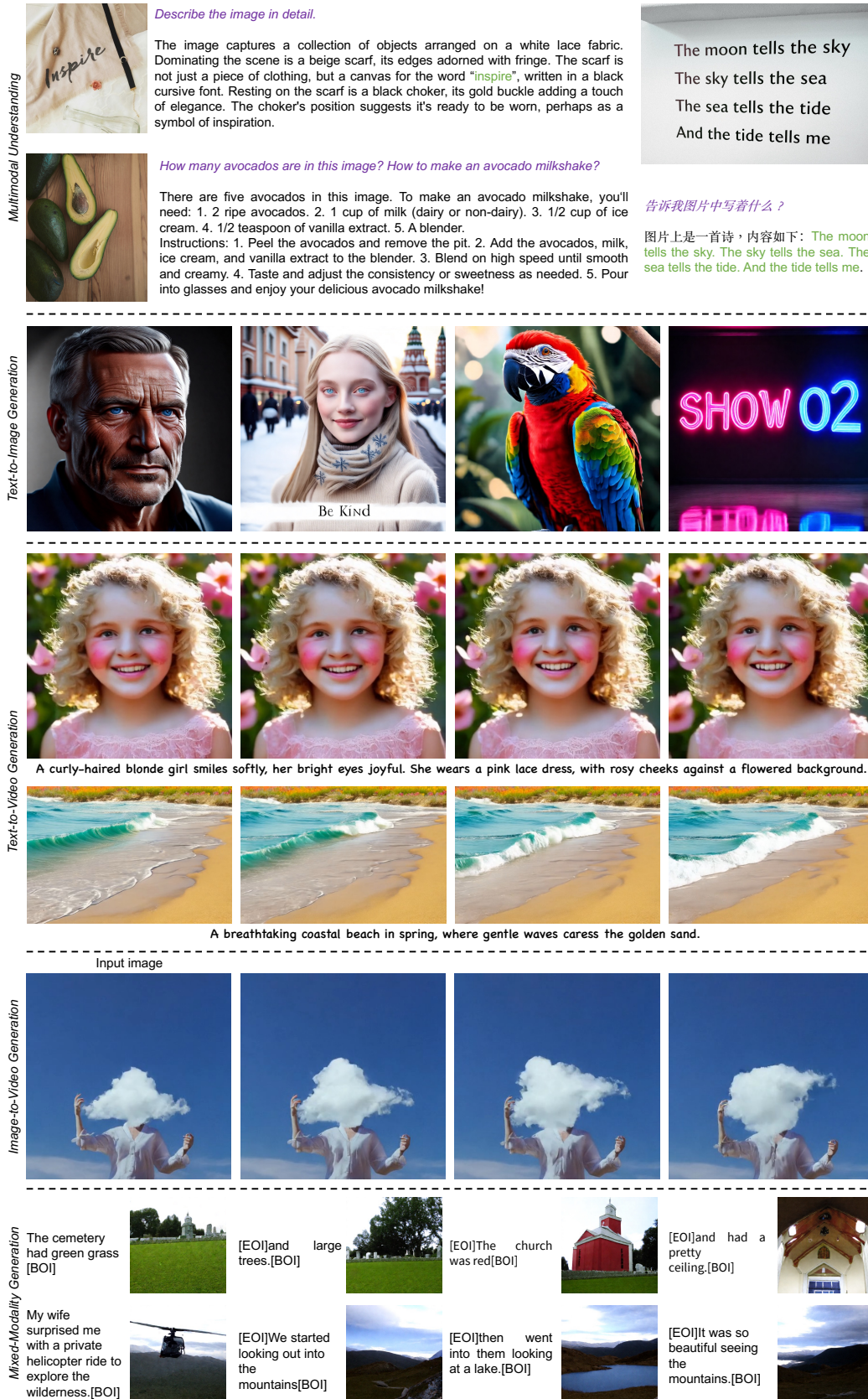


Figure 2: Multimodal understanding and generation examples.

Table 8: Comparison with text-to-video models on the VBench [48] benchmark. # Params. indicates the number of total parameters for video generation including the base LLM and flow head. QS: Quality Score, SS: Semantic Score, SC: Subject Consistency, BC: Background Consistency, TF: Temporal Flickering, MS: Motion Smoothness, DD: Dynamic Degree, AQ: Aesthetic Quality, IQ: Imaging Quality, OC: Object Class, MO: Multiple Objects, HA: Human Action, C: Color, SR: Spatial Relationship, S: Scene, AS: Appearance style, TS: Temporal Style, OC’: Overall Consistency.

Models	# Params.	Total	QS	SS	SC	BC	TF	MS	DD	AQ	IQ	OC	MO	HA	C	SR	S	AS	TS	OC’
ModelScope [114]	1.7B	75.75	78.05	66.54	89.87	95.29	98.28	95.79	66.39	52.06	58.57	82.25	38.98	92.40	81.72	33.68	39.26	23.39	25.37	25.67
LaVie [116]	3B	77.08	78.78	70.31	91.41	97.47	98.30	96.38	49.72	54.94	61.90	91.82	33.32	96.80	86.39	34.09	52.69	23.56	25.93	26.41
OpenSoraPlan V1.3 [67]	-	77.23	80.14	65.62	97.79	97.24	99.20	99.05	30.28	60.42	56.21	85.56	43.58	86.80	79.30	51.61	36.73	20.03	22.47	24.47
Show-1 [141]	6B	78.93	80.42	72.98	95.53	98.02	99.12	98.24	44.44	57.35	58.66	93.07	45.47	95.60	86.35	53.50	47.03	23.06	25.28	27.46
AnimateDiff-V2 [43]	-	80.27	82.90	69.75	95.30	97.68	98.75	97.76	40.83	67.16	70.10	90.90	36.88	92.60	87.47	34.60	50.19	22.42	26.03	27.04
Gen-2 [1]	-	80.58	82.47	73.03	97.61	97.61	99.56	99.58	18.89	66.96	67.42	90.92	55.47	89.20	89.49	66.91	48.91	19.34	24.12	26.17
Pika-1.0 [2]	-	80.69	82.92	71.77	96.94	97.36	99.74	99.50	47.50	62.04	61.87	88.72	43.08	86.20	90.57	61.03	49.83	22.26	24.22	25.94
VideoCrafter-2.0 [16]	-	80.44	82.20	73.42	96.85	98.22	98.41	97.73	42.50	63.13	67.22	92.55	40.66	95.00	92.92	35.86	55.29	25.13	25.84	28.23
CogVideoX [137]	5B	81.61	82.75	77.04	96.23	96.52	98.66	96.92	70.97	61.98	62.90	85.23	62.11	99.40	82.81	66.35	53.20	24.91	25.38	27.59
Kling [4]	-	81.85	83.39	75.68	98.33	97.60	99.30	99.40	46.94	61.21	65.62	87.24	68.05	93.40	89.90	73.03	50.86	19.62	24.17	26.42
Step-Video-T2V [79]	30B	81.83	84.46	71.28	98.05	97.67	99.40	99.08	53.06	61.23	70.63	80.56	50.55	94.00	88.25	71.47	24.38	23.17	26.01	27.12
Gen-3 [3]	-	82.32	84.11	75.17	97.10	96.62	98.61	99.23	60.14	63.34	66.82	87.81	53.64	96.40	80.90	65.09	54.57	24.31	24.71	26.69
Emu3 [115]	8B	80.96	-	-	95.32	97.69	-	98.93	79.27	59.64	-	86.17	44.64	77.71	-	68.73	37.11	20.92	-	-
VILA-U [124]	7B	74.01	76.26	65.04	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
HaploOmni [125]	9B	78.10	-	-	96.40	97.60	-	96.80	65.30	-	-	-	-	-	-	-	34.60	-	-	-
Show-o2 (Ours)	2B	81.34	82.10	78.31	97.28	96.78	97.68	98.25	40.83	65.15	67.06	94.81	76.01	95.20	80.89	62.61	57.67	23.29	25.27	27.00

Table 9: Comparison with image-to-video models on the VBench [48] benchmark.

Models	I2V Subject	I2V Background	Camera Motion	Subject Consistency	Background Consistency	Temporal Flickering	Motion Smoothness	Dynamic Degree	Aesthetic Quality	Imaging Quality
DynamiCrafter-1024 [131]	96.71	96.05	35.44	95.69	97.38	97.63	97.38	47.40	66.46	69.34
SEINE-512x320 [25]	94.85	94.02	23.36	94.20	97.26	96.72	96.68	34.31	58.42	70.97
I2VGen-XL [145]	96.74	95.44	13.32	96.36	97.93	98.48	98.31	24.96	65.33	69.85
Animate-Anything [31]	98.54	96.88	12.56	98.90	98.19	98.14	98.61	2.68	67.12	72.09
ConsistI2V [92]	94.69	94.57	33.60	95.27	98.28	97.56	97.38	18.62	59.00	66.92
VideoCrafter-I2V [15]	90.97	90.51	33.58	97.86	98.79	98.19	98.00	22.60	60.78	71.68
SVD-XT-1.1 [11]	97.51	97.62	-	95.42	96.77	99.17	98.12	43.17	60.23	70.23
MarDini [74]	98.78	96.46	-	-	-	-	-	-	-	-
Show-o2 (Ours)	96.94	98.83	28.41	93.83	97.45	-	97.76	25.85	61.92	69.87

4.3 Mixed-Modality Generation

We demonstrate mixed-modality generation capabilities of our model using downstream task visual storytelling dataset [47] in Fig. 2. During fine-tuning, given an interleaved image-text sequence, we apply noise to all images in the sequence with a probability of 0.3. Otherwise, we randomly retain a number of the earlier images in the sequence and only apply noise to the later ones. Benefiting from the general interleaved sequence format mentioned in 3.1, our model can predict the [BOI] once it begins to generate an image. Upon detecting the [BOI] token, noises will be appended to the sequence to gradually generate an image. The generated text tokens and images will be served as context to continue generating the following output. Fig. 2 includes two examples demonstrating our model’s ability to interleavely generate coherent text and images, vividly narrating a story.

4.4 Ablation Studies

We show the pilot study results in Table 10, which validated the effect of spatial (-temporal) fusion on multimodal understanding and generation performance. For efficiency, we adopt LLaMA-3.2-1B as the base language model and use only around 1M multimodal understanding data and the ImageNet-1K generation data [33]. Under the same training settings, there are improvements in terms of both multimodal understanding and generation metrics, including MME-p, GQA, and FID-5K. This validates that the involved semantic and low-level features in the fusion mechanism would potentially help both the multimodal generation and understanding capabilities to some extent.

Table 10: Impact of spatial (-temporal) fusion.

	MME-p ↑	GQA ↑	POPE ↑	FID-5K ↓
w/o Fusion	1164.7	56.2	82.6	21.8
w Fusion	1187.8	57.6	82.6	20.5

Table 11 provides the effect of training stages on the generation performance on the GenEval and DPG-Bench benchmarks. One can observe that stage-2 training consistently and significantly improves both metrics, which validates the importance of the second stage.

Table 11: Effect of training stages.

Stage-1	Stage-2	GenEval	DPG-Bench
✓		0.63	83.28
✓	✓	0.73	84.70

We perform ablation studies to examine the effect of classifier-free guidance (CFG) and inference steps on the performance using the 1.5B model. As shown in Table 12, increasing the CFG guidance scale and inference steps (in a range) would potentially improve the GenEval and DPG-Bench scores. However, the improvements of the GenEval score are not significant when the CFG guidance is set as larger than 5.0.

Table 12: Effect of CFG guidance and inference steps.

CFG guidance	Inference steps	GenEval	DPG-Bench
2.5	50	0.65	81.6
5.0	50	0.71	83.9
7.5	50	0.71	84.8
10	50	0.71	85.0
7.5	25	0.71	84.6
7.5	100	0.73	84.7

Table 13: Impact of training recipe on text-only performance. One-stage training denotes full-parameter-co-training on image-text pairs and the text-only RefinedWeb [88] data. Note that the curated multimodal understanding data consists of text-only instructional data. We perform the evaluation under the same setting using the [lm-evaluation-harness tool](#).

Models	# Params.	Training Recipe	MMLU	GPQA	GSM8K	HumanEval
Qwen2.5 Instruct [134]	1.5B	-	60.20 $\pm$ 0.39	28.12 $\pm$ 2.13	51.86 $\pm$ 1.38	35.37 $\pm$ 3.74
Show-o2 (Ours)	1.5B	One-stage training with RefinedWeb	28.25 $\pm$ 0.38	25.00 $\pm$ 2.05	4.55 $\pm$ 0.57	3.05 $\pm$ 1.35
Show-o2 (Ours)	1.5B	Our two-stage training	56.75 $\pm$ 1.37	29.24 $\pm$ 2.15	49.43 $\pm$ 1.38	35.54 $\pm$ 3.70
Qwen2.5 Instruct [134]	7B	-	71.75 $\pm$ 0.36	32.37 $\pm$ 2.21	82.49 $\pm$ 1.05	65.24 $\pm$ 3.73
Show-o2 (Ours)	7B	One-stage training with RefinedWeb	28.43 $\pm$ 0.21	26.34 $\pm$ 2.08	1.52 $\pm$ 0.34	4.01 $\pm$ 1.25
Show-o2 (Ours)	7B	Our two-stage training	70.73 $\pm$ 0.36	31.47 $\pm$ 2.22	75.28 $\pm$ 1.19	70.73 $\pm$ 3.56

Table 13 shows that our models effectively preserve language knowledge and achieve performance comparable to the original Qwen2.5-1.5B and Qwen2.5-7B Instruct models. In contrast, direct one-stage full-parameter-co-training with textual data such as RefinedWeb results in substantial performance degradation, highlighting the necessity of the two-stage training approach when high-quality corpora are unavailable.

As shown in Table 14, our ablation study reveals that increasing the number of image tokens significantly boosts performance across all tasks, even though the model was trained with a fixed image resolution. Using the AnyRes strategy at inference time consistently improves results, highlighting the benefit of higher token counts for capturing fine-grained details. When compared to the baseline LLaVA-OV-7B, our model achieves comparable results on DocVQA, InfoVQA, and TextVQA validation sets, but underperforms on ChartQA. We attribute this gap to the limited chart-related data available during semantic layer distillation, which constrains the model’s ability to capture chart-specific information. We believe that incorporating more OCR and document-centric data into the distillation process will further strengthen the unified model’s OCR and document understanding capabilities.

Table 14: Impact of image token count on chart, text, and document VQA.

Models	# Params.	# Image tokens	ChartQA	DocVQA _{val}	InfoVQA _{val}	TextVQA _{val}
LLaVA-OV	7B	729	56.24	62.71	39.59	66.19
Show-o2	7B	729	48.00	59.34	42.31	62.92
Show-o2	7B	5 $\times$ 729	66.92	77.26	45.80	71.54

5 Conclusion

This paper proposed native unified multimodal models, *i.e.*, Show-o2, scalable for multimodal understanding and generation, image and video modalities, by integrating 3D causal VAE, autoregressive modeling, and flow matching. A dual-path of spatial (-temporal) fusion mechanism guided the construction of unified visual representations with both high- and low-level features. A two-stage training recipe enables effective learning of unified capabilities, resulting in a versatile model capable of handling diverse tasks, including multimodal understanding and image/video generation. Extensive experiments demonstrate the model’s state-of-the-art performance across various benchmarks.

Ethics Statement. We use CC12M [14], COYO [12], LAION-Aesthetic-12M, AI-synthetic data, and internally collected video data, all of which are strictly filtered to exclude copyrighted materials and visible watermarks during data curation. A safety checker is integrated into our inference pipeline to mitigate potential misuse. We have conducted repeated evaluations and observed stable results with minimal variations. However, our experimental comparisons follow prior literature that did not report error bars. For language-only evaluations, we report the standard deviations.

Acknowledgments. This research is supported by the National Research Foundation, Singapore under its AI Singapore Programme (AISG Award No: AISG3-RP-2022-030).

References

- [1] Gen-2. Accessed September 25, 2023 [Online] <https://research.runwayml.com/gen2>, 2023.
- [2] Pika 1.0. Accessed December 28, 2023 [Online] <https://www.pika.art/>, 2023.
- [3] Gen-3. Accessed June 17, 2024 [Online] <https://runwayml.com/research/introducing-gen-3-alpha>, 2024.
- [4] Kling. Accessed June 6, 2024 [Online] <https://klingai.kuaishou.com/>, 2024.
- [5] Inclusion AI. Ming-omni: A unified multimodal model for perception and generation, 2025.
- [6] Jinze Bai, Shuai Bai, Shusheng Yang, Shijie Wang, Sinan Tan, Peng Wang, Junyang Lin, Chang Zhou, and Jingren Zhou. Qwen-vl: A versatile vision-language model for understanding, localization, text reading, and beyond. *arXiv preprint arXiv:2308.12966*, 2023.
- [7] Shuai Bai, Keqin Chen, Xuejing Liu, Jialin Wang, Wenbin Ge, Sibao Song, Kai Dang, Peng Wang, Shijie Wang, Jun Tang, Humen Zhong, Yuanzhi Zhu, Mingkun Yang, Zhaohai Li, Jianqiang Wan, Pengfei Wang, Wei Ding, Zheren Fu, Yiheng Xu, Jiabo Ye, Xi Zhang, Tianbao Xie, Zesen Cheng, Hang Zhang, Zhibo Yang, Haiyang Xu, and Junyang Lin. Qwen2.5-vl technical report. *arXiv preprint arXiv:2502.13923*, 2025.
- [8] Max Bain, Arsha Nagrani, Gül Varol, and Andrew Zisserman. Frozen in time: A joint video and image encoder for end-to-end retrieval. In *IEEE International Conference on Computer Vision*, 2021.
- [9] Fan Bao, Shen Nie, Kaiwen Xue, Yue Cao, Chongxuan Li, Hang Su, and Jun Zhu. All are worth words: A vit backbone for diffusion models. In *CVPR*, 2023.
- [10] James Betker, Gabriel Goh, Li Jing, † TimBrooks, Jianfeng Wang, Linjie Li, † LongOuyang, † Jun-tangZhuang, † JoyceLee, † YufeiGuo, † WesamManassra, † PrafullaDhariwal, † CaseyChu, † YunxinJiao, and Aditya Ramesh. Improving image generation with better captions.
- [11] 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.
- [12] Minwoo Byeon, Beomhee Park, Haechon Kim, Sungjun Lee, Woonhyuk Baek, and Saehoon Kim. Coyo-700m: Image-text pair dataset. <https://github.com/kakaobrain/coyo-dataset>, 2022.
- [13] Jingjing Chang, Yixiao Fang, Peng Xing, Shuhan Wu, Wei Cheng, Rui Wang, Xianfang Zeng, Gang Yu, and Hai-Bao Chen. Oneig-bench: Omni-dimensional nuanced evaluation for image generation. *arXiv preprint arxiv:2506.07977*, 2025.
- [14] Soravit Changpinyo, Piyush Sharma, Nan Ding, and Radu Soricut. Conceptual 12m: Pushing web-scale image-text pre-training to recognize long-tail visual concepts. In *CVPR*, pages 3558–3568, 2021.
- [15] Haoxin Chen, Menghan Xia, Yingqing He, Yong Zhang, Xiaodong Cun, Shaoshu Yang, Jinbo Xing, Yaofang Liu, Qifeng Chen, Xintao Wang, Chao Weng, and Ying Shan. Videocrafter1: Open diffusion models for high-quality video generation. *arXiv preprint arXiv:2310.19512*, 2023.
- [16] Haoxin Chen, Yong Zhang, Xiaodong Cun, Menghan Xia, Xintao Wang, Chao Weng, and Ying Shan. Videocrafter2: Overcoming data limitations for high-quality video diffusion models, 2024.
- [17] Junsong Chen, Chongjian Ge, Enze Xie, Yue Wu, Lewei Yao, Xiaozhe Ren, Zhongdao Wang, Ping Luo, Huchuan Lu, and Zhenguo Li. Pixart-sigma: Weak-to-strong training of diffusion transformer for 4k text-to-image generation, 2024.
- [18] Jiuhai Chen, Zhiyang Xu, Xichen Pan, Yushi Hu, Can Qin, Tom Goldstein, Lifu Huang, Tianyi Zhou, Saining Xie, Silvio Savarese, et al. Blip3-o: A family of fully open unified multimodal models-architecture, training and dataset. *arXiv preprint arXiv:2505.09568*, 2025.
- [19] Junsong Chen, Jincheng Yu, Chongjian Ge, Lewei Yao, Enze Xie, Zhongdao Wang, James T. Kwok, Ping Luo, Huchuan Lu, and Zhenguo Li. Pixart- α : Fast training of diffusion transformer for photorealistic text-to-image synthesis. In *ICLR*. OpenReview.net, 2024.
- [20] Lin Chen, Jisong Li, Xiaoyi Dong, Pan Zhang, Conghui He, Jiaqi Wang, Feng Zhao, and Dahua Lin. Sharegpt4v: Improving large multi-modal models with better captions. *arXiv preprint arXiv:2311.12793*, 2023.
- [21] Lin Chen, Jinsong Li, Xiaoyi Dong, Pan Zhang, Yuhang Zang, Zehui Chen, Haodong Duan, Jiaqi Wang, Yu Qiao, Dahua Lin, et al. Are we on the right way for evaluating large vision-language models? *arXiv preprint arXiv:2403.20330*, 2024.
- [22] Mark Chen, Alec Radford, Rewon Child, Jeffrey Wu, Heewoo Jun, David Luan, and Ilya Sutskever. Generative pretraining from pixels. In *ICML*, pages 1691–1703, 2020.
- [23] Tsai-Shien Chen, Aliaksandr Siarohin, Willi Menapace, Ekaterina Deyneka, Hsiang-wei Chao, Byung Eun Jeon, Yuwei Fang, Hsin-Ying Lee, Jian Ren, Ming-Hsuan Yang, and Sergey Tulyakov. Panda-70m: Captioning 70m videos with multiple cross-modality teachers. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2024.
- [24] Wei Chen, Lin Li, Yongqi Yang, Bin Wen, Fan Yang, Tingting Gao, Yu Wu, and Long Chen. Comm: A coherent interleaved image-text dataset for multimodal understanding and generation. *arXiv preprint arXiv:2406.10462*, 2024.

- [25] Xinyuan Chen, Yaohui Wang, Lingjun Zhang, Shaobin Zhuang, Xin Ma, Jiashuo Yu, Yali Wang, Dahua Lin, Yu Qiao, and Ziwei Liu. Seine: Short-to-long video diffusion model for generative transition and prediction. *arXiv preprint arXiv:2310.20700*, 2023.
- [26] 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.
- [27] 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.
- [28] Zisheng Chen, Chunwei Wang, Xiuwei Chen, Hang Xu, Jianhua Han, and Xiaodan Liang. Semhitok: A unified image tokenizer via semantic-guided hierarchical codebook for multimodal understanding and generation. *arXiv preprint arXiv:2503.06764*, 2025.
- [29] Zhe Chen, Weiyun Wang, Yue Cao, Yangzhou Liu, Zhangwei Gao, Erfei Cui, Jinguo Zhu, Shenglong Ye, Hao Tian, Zhaoyang Liu, et al. Expanding performance boundaries of open-source multimodal models with model, data, and test-time scaling. *arXiv preprint arXiv:2412.05271*, 2024.
- [30] Zesen Cheng, Sicong Leng, Hang Zhang, Yifei Xin, Xin Li, Guanzheng Chen, Yongxin Zhu, Wenqi Zhang, Ziyang Luo, Deli Zhao, and Lidong Bing. Videollama 2: Advancing spatial-temporal modeling and audio understanding in video-llms, 2024.
- [31] Zuozhuo Dai, Zhenghao Zhang, Yao Yao, Bingxue Qiu, Siyu Zhu, Long Qin, and Weizhi Wang. Fine-grained open domain image animation with motion guidance. *arXiv preprint arXiv:2311.12886*, 2023.
- [32] Chaorui Deng, Deyao Zhu, Kunchang Li, Chenhui Gou, Feng Li, Zeyu Wang, Shu Zhong, Weihao Yu, Xiaonan Nie, Ziang Song, et al. Emerging properties in unified multimodal pretraining. *arXiv preprint arXiv:2505.14683*, 2025.
- [33] Jia Deng, Wei Dong, Richard Socher, Li-Jia Li, Kai Li, and Li Fei-Fei. Imagenet: A large-scale hierarchical image database. In *CVPR*, pages 248–255, 2009.
- [34] Haiwen Diao, Yufeng Cui, Xiaotong Li, Yueze Wang, Huchuan Lu, and Xinlong Wang. Unveiling encoder-free vision-language models. *arXiv preprint arXiv:2406.11832*, 2024.
- [35] Haiwen Diao, Xiaotong Li, Yufeng Cui, Yueze Wang, Haoge Deng, Ting Pan, Wenxuan Wang, Huchuan Lu, and Xinlong Wang. Evev2: Improved baselines for encoder-free vision-language models. *arXiv preprint arXiv:2502.06788*, 2025.
- [36] Runpei Dong, Chunrui Han, Yuang Peng, Zekun Qi, Zheng Ge, Jinrong Yang, Liang Zhao, Jianjian Sun, Hongyu Zhou, Haoran Wei, Xiangwen Kong, Xiangyu Zhang, Kaisheng Ma, and Li Yi. DreamLLM: Synergistic multimodal comprehension and creation. In *ICLR*, 2024.
- [37] Patrick Esser, Sumith Kulal, Andreas Blattmann, Rahim Entezari, Jonas Müller, Harry Saini, Yam Levi, Dominik Lorenz, Axel Sauer, Frederic Boesel, et al. Scaling rectified flow transformers for high-resolution image synthesis. In *ICML*, 2024.
- [38] Lijie Fan, Luming Tang, Siyang Qin, Tianhong Li, Xuan Yang, Siyuan Qiao, Andreas Steiner, Chen Sun, Yuanzhen Li, Tao Zhu, et al. Unified autoregressive visual generation and understanding with continuous tokens. *arXiv preprint arXiv:2503.13436*, 2025.
- [39] Chaoyou Fu, Peixian Chen, Yunhang Shen, Yulei Qin, Mengdan Zhang, Xu Lin, Zhenyu Qiu, Wei Lin, Jinrui Yang, Xiawu Zheng, Ke Li, Xing Sun, and Rongrong Ji. MME: A comprehensive evaluation benchmark for multimodal large language models. *CoRR*, abs/2306.13394, 2023.
- [40] Yuying Ge, Sijie Zhao, Jinguo Zhu, Yixiao Ge, Kun Yi, Lin Song, Chen Li, Xiaohan Ding, and Ying Shan. Seed-x: Multimodal models with unified multi-granularity comprehension and generation. *arXiv preprint arXiv:2404.14396*, 2024.
- [41] Dhruva Ghosh, Hannaneh Hajishirzi, and Ludwig Schmidt. Geneval: An object-focused framework for evaluating text-to-image alignment. In *NeurIPS*, 2023.
- [42] Biao Gong, Cheng Zou, Dandan Zheng, Hu Yu, Jingdong Chen, Jianxin Sun, Junbo Zhao, Jun Zhou, Kaixiang Ji, Lixiang Ru, et al. Ming-lite-uni: Advancements in unified architecture for natural multimodal interaction. *arXiv preprint arXiv:2505.02471*, 2025.
- [43] Yuwei Guo, Ceyuan Yang, Anyi Rao, Zhengyang Liang, Yaohui Wang, Yu Qiao, Maneesh Agrawala, Dahua Lin, and Bo Dai. Animatediff: Animate your personalized text-to-image diffusion models without specific tuning. *International Conference on Learning Representations*, 2024.
- [44] HiDream-ai. Hidream-1l. <https://github.com/HiDream-ai/HiDream-1l>, 2025.
- [45] Xiwei Hu, Rui Wang, Yixiao Fang, Bin Fu, Pei Cheng, and Gang Yu. Ella: Equip diffusion models with llm for enhanced semantic alignment, 2024.
- [46] Runhui Huang, Chunwei Wang, Junwei Yang, Guansong Lu, Yunlong Yuan, Jianhua Han, Lu Hou, Wei Zhang, Lanqing Hong, Hengshuang Zhao, and Hang Xu. Illume+: Illuminating unified mllm with dual visual tokenization and diffusion refinement. *arXiv preprint arXiv:2504.01934*, 2025.
- [47] Ting-Hao K. Huang, Francis Ferraro, Nasrin Mostafazadeh, Ishan Misra, Jacob Devlin, Aishwarya Agrawal, Ross Girshick, Xiaodong He, Pushmeet Kohli, Dhruv Batra, et al. Visual storytelling. In *15th Annual Conference of the North American Chapter of the Association for Computational Linguistics (NAACL 2016)*, 2016.

- [48] Ziqi Huang, Yanan He, Jiashuo Yu, Fan Zhang, Chenyang Si, Yuming Jiang, Yuanhan Zhang, Tianxing Wu, Qingyang Jin, Nattapol Chanpaisit, Yaohui Wang, Xinyuan Chen, Limin Wang, Dahua Lin, Yu Qiao, and Ziwei Liu. VBench: Comprehensive benchmark suite for video generative models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2024.
- [49] Drew A. Hudson and Christopher D. Manning. GQA: A new dataset for real-world visual reasoning and compositional question answering. In *CVPR*, pages 6700–6709. Computer Vision Foundation / IEEE, 2019.
- [50] Yang Jiao, Haibo Qiu, Zequn Jie, Shaoxiang Chen, Jingjing Chen, Lin Ma, and Yu-Gang Jiang. Unitoken: Harmonizing multimodal understanding and generation through unified visual encoding. *arXiv preprint arXiv:2504.04423*, 2025.
- [51] Aniruddha Kembhavi, Mike Salvato, Eric Kolve, Minjoon Seo, Hannaneh Hajishirzi, and Ali Farhadi. A diagram is worth a dozen images. In *Computer Vision—ECCV 2016: 14th European Conference, Amsterdam, The Netherlands, October 11–14, 2016, Proceedings, Part IV 14*, pages 235–251. Springer, 2016.
- [52] Dan Kondratyuk, Lijun Yu, Xiuye Gu, José Lezama, Jonathan Huang, Rachel Hornung, Hartwig Adam, Hassan Akbari, Yair Alon, Vighnesh Birodkar, et al. Videopoet: A large language model for zero-shot video generation. *arXiv preprint arXiv:2312.14125*, 2023.
- [53] Siqi Kou, Jiachun Jin, Chang Liu, Ye Ma, Jian Jia, Quan Chen, Peng Jiang, and Zhijie Deng. Orthus: Autoregressive interleaved image-text generation with modality-specific heads. *arXiv preprint arXiv:2412.00127*, 2024.
- [54] Black Forest Labs. Flux. <https://github.com/black-forest-labs/flux>, 2024.
- [55] Bohao Li, Rui Wang, Guangzhi Wang, Yuying Ge, Yixiao Ge, and Ying Shan. Seed-bench: Benchmarking multimodal llms with generative comprehension. *arXiv preprint arXiv:2307.16125*, 2023.
- [56] Bo Li, Yuanhan Zhang, Dong Guo, Renrui Zhang, Feng Li, Hao Zhang, Kaichen Zhang, Yanwei Li, Ziwei Liu, and Chunyuan Li. Llava-onevision: Easy visual task transfer. *arXiv preprint arXiv:2408.03326*, 2024.
- [57] Daiqing Li, Aleks Kamko, Ehsan Akhgari, Ali Sabet, Linmiao Xu, and Suhail Doshi. Playground v2.5: Three insights towards enhancing aesthetic quality in text-to-image generation. *ArXiv*, abs/2402.17245, 2024.
- [58] Hao Li, Changyao Tian, Jie Shao, Xizhou Zhu, Zhaokai Wang, Jinguo Zhu, Wenhan Dou, Xiaogang Wang, Hongsheng Li, Lewei Lu, and Jifeng Dai. Synergen-vl: Towards synergistic image understanding and generation with vision experts and token folding. *arXiv preprint arXiv:2412.09604*, 2024.
- [59] Haopeng Li, Jinyue Yang, Guoqi Li, and Huan Wang. Autoregressive image generation with randomized parallel decoding. *arXiv preprint arXiv:2503.10568*, 2025.
- [60] Qingyun Li, Zhe Chen, Weiyun Wang, Wenhai Wang, Shenglong Ye, Zhenjiang Jin, et al. Omnicorpus: A unified multimodal corpus of 10 billion-level images interleaved with text. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [61] Shufan Li, Konstantinos Kallidromitis, Hritik Bansal, Akash Gokul, Yusuke Kato, Kazuki Kozuka, Jason Kuen, Zhe Lin, Kai-Wei Chang, and Aditya Grover. Lavidia: A large diffusion language model for multimodal understanding. *arXiv preprint arXiv:2505.16839*, 2025.
- [62] Tianhong Li, Yonglong Tian, He Li, Mingyang Deng, and Kaiming He. Autoregressive image generation without vector quantization. *arXiv preprint arXiv:2406.11838*, 2024.
- [63] Xiaotong Li, Fan Zhang, Haiwen Diao, Yueze Wang, Xinlong Wang, and Ling-Yu Duan. Densefusion-1m: Merging vision experts for comprehensive multimodal perception. *2407.08303*, 2024.
- [64] Zijie Li, Henry Li, Yichun Shi, Amir Barati Farimani, Yuval Kluger, Linjie Yang, and Peng Wang. Dual diffusion for unified image generation and understanding. *arXiv preprint arXiv:2501.00289*, 2024.
- [65] Zhimin Li, Jianwei Zhang, Qin Lin, Jiangfeng Xiong, Yanxin Long, Xincheng Deng, Yingfang Zhang, Xingchao Liu, Minbin Huang, Zedong Xiao, Dayou Chen, Jiajun He, Jiahao Li, Wenyue Li, Chen Zhang, Rongwei Quan, Jianxiang Lu, Jiabin Huang, Xiaoyan Yuan, Xiaoxiao Zheng, Yixuan Li, Jihong Zhang, Chao Zhang, Meng Chen, Jie Liu, Zheng Fang, Weiyan Wang, Jinbao Xue, Yangyu Tao, Jianchen Zhu, Kai Liu, Sihuan Lin, Yifu Sun, Yun Li, Dongdong Wang, Mingtao Chen, Zhichao Hu, Xiao Xiao, Yan Chen, Yuhong Liu, Wei Liu, Di Wang, Yong Yang, Jie Jiang, and Qinglin Lu. Hunyuan-dit: A powerful multi-resolution diffusion transformer with fine-grained chinese understanding, 2024.
- [66] Chao Liao, Liyang Liu, Xun Wang, Zhengxiong Luo, Xinyu Zhang, Wenliang Zhao, Jie Wu, Liang Li, Zhi Tian, and Weilin Huang. Mogao: An omni foundation model for interleaved multi-modal generation. *arXiv preprint arXiv:2505.05472*, 2025.
- [67] Bin Lin, Yunyang Ge, Xinhua Cheng, Zongjian Li, Bin Zhu, Shaodong Wang, Xianyi He, Yang Ye, Shenghai Yuan, Liuhan Chen, et al. Open-sora plan: Open-source large video generation model. *arXiv preprint arXiv:2412.00131*, 2024.
- [68] Bin Lin, Zongjian Li, Xinhua Cheng, Yuwei Niu, Yang Ye, Xianyi He, Shenghai Yuan, Wangbo Yu, Shaodong Wang, Yunyang Ge, et al. Uniworld: High-resolution semantic encoders for unified visual understanding and generation. *arXiv preprint arXiv:2506.03147*, 2025.

- [69] Haokun Lin, Teng Wang, Yixiao Ge, Yuying Ge, Zhichao Lu, Ying Wei, Qingfu Zhang, Zhenan Sun, and Ying Shan. Toklip: Marry visual tokens to clip for multimodal comprehension and generation. *arXiv preprint arXiv:2505.05422*, 2025.
- [70] Ji Lin, Hongxu Yin, Wei Ping, Pavlo Molchanov, Mohammad Shoybi, and Song Han. Vila: On pre-training for visual language models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 26689–26699, 2024.
- [71] Yaron Lipman, Ricky T. Q. Chen, Heli Ben-Hamu, Maximilian Nickel, and Matthew Le. Flow matching for generative modeling. In *The Eleventh International Conference on Learning Representations*, 2023.
- [72] Haotian Liu, Chunyuan Li, Yuheng Li, and Yong Jae Lee. Improved baselines with visual instruction tuning. In *CVPR*, pages 26296–26306, 2024.
- [73] Haotian Liu, Chunyuan Li, Qingyang Wu, and Yong Jae Lee. Visual instruction tuning. *NeurIPS*, 36, 2024.
- [74] Haozhe Liu, Shikun Liu, Zijian Zhou, Mengmeng Xu, Yanping Xie, Xiao Han, Juan C. Pérez, Ding Liu, Kumara Kahatapitiya, Menglin Jia, Jui-Chieh Wu, Sen He, Tao Xiang, Jürgen Schmidhuber, and Juan-Manuel Pérez-Rúa. Mardini: Masked autoregressive diffusion for video generation at scale. *arXiv preprint arXiv:2410.20280*, 2024.
- [75] Xingchao Liu, Chengyue Gong, and Qiang Liu. Flow straight and fast: Learning to generate and transfer data with rectified flow. *arXiv preprint arXiv:2209.03003*, 2022.
- [76] Yuan Liu, Haodong Duan, Yuanhan Zhang, Bo Li, Songyang Zhang, Wangbo Zhao, Yike Yuan, Jiaqi Wang, Conghui He, Ziwei Liu, et al. Mmbench: Is your multi-modal model an all-around player? In *European conference on computer vision*, pages 216–233. Springer, 2024.
- [77] Jiasen Lu, Christopher Clark, Sangho Lee, Zichen Zhang, Savva Khosla, Ryan Marten, Derek Hoiem, and Aniruddha Kembhavi. Unified-io 2: Scaling autoregressive multimodal models with vision, language, audio, and action. *arXiv preprint arXiv:2312.17172*, 2023.
- [78] Chuofan Ma, Yi Jiang, Junfeng Wu, Jihan Yang, Xin Yu, Zehuan Yuan, Bingyue Peng, and Xiaojuan Qi. Unitok: A unified tokenizer for visual generation and understanding. *arXiv preprint arXiv:2502.20321*, 2025.
- [79] Guoqing Ma, Haoyang Huang, Kun Yan, Liangyu Chen, Nan Duan, Shengming Yin, Changyi Wan, Ranchen Ming, Xiaoniu Song, Xing Chen, Yu Zhou, Deshan Sun, Deyu Zhou, Jian Zhou, Kaijun Tan, Kang An, Mei Chen, Wei Ji, Qiling Wu, Wen Sun, Xin Han, Yanan Wei, Zheng Ge, Aojie Li, Bin Wang, Bizhu Huang, Bo Wang, Brian Li, Changxing Miao, Chen Xu, Chenfei Wu, Chenguang Yu, Dapeng Shi, Dingyuan Hu, Enle Liu, Gang Yu, Ge Yang, Guanzhe Huang, Gulin Yan, Haiyang Feng, Hao Nie, Haonan Jia, Hanpeng Hu, Hanqi Chen, Haolong Yan, Heng Wang, Hongcheng Guo, Huilin Xiong, Huixin Xiong, Jiahao Gong, Jianchang Wu, Jiaoren Wu, Jie Wu, Jie Yang, Jiashuai Liu, Jiashuo Li, Jingyang Zhang, Junjing Guo, Junzhe Lin, Kaixiang Li, Lei Liu, Lei Xia, Liang Zhao, Liguang Tan, Liwen Huang, Liying Shi, Ming Li, Mingliang Li, Muhua Cheng, Na Wang, Qiaohui Chen, Qinglin He, Qiuyan Liang, Quan Sun, Ran Sun, Rui Wang, Shaoliang Pang, Shiliang Yang, Sitong Liu, Siqi Liu, Shuli Gao, Tiancheng Cao, Tianyu Wang, Weipeng Ming, Wenqing He, Xu Zhao, Xuelin Zhang, Xianfang Zeng, Xiaojia Liu, Xuan Yang, Yaqi Dai, Yanbo Yu, Yang Li, Yineng Deng, Yingming Wang, Yilei Wang, Yuanwei Lu, Yu Chen, Yu Luo, Yuchu Luo, Yuhe Yin, Yuheng Feng, Yuxiang Yang, Zecheng Tang, Zekai Zhang, Zidong Yang, Binxing Jiao, Jiansheng Chen, Jing Li, Shuchang Zhou, Xiangyu Zhang, Xinhao Zhang, Yibo Zhu, Heung-Yeung Shum, and Daxin Jiang. Step-video-t2v technical report: The practice, challenges, and future of video foundation model, 2025.
- [80] Yiyang Ma, Xingchao Liu, Xiaokang Chen, Wen Liu, Chengyue Wu, Zhiyu Wu, Zizheng Pan, Zhenda Xie, Haowei Zhang, Xingkai yu, Liang Zhao, Yisong Wang, Jiaying Liu, and Chong Ruan. Janusflow: Harmonizing autoregression and rectified flow for unified multimodal understanding and generation, 2024.
- [81] Kepan Nan, Rui Xie, Penghao Zhou, Tiehan Fan, Zhenheng Yang, Zhijie Chen, Xiang Li, Jian Yang, and Ying Tai. Openvid-1m: A large-scale high-quality dataset for text-to-video generation. *arXiv preprint arXiv:2407.02371*, 2024.
- [82] OpenAI. Gpt-4v. <https://openai.com/index/gpt-4v-system-card/>, 2023.
- [83] OpenAI. Hello gpt-4o. <https://openai.com/index/hello-gpt-4o/>, 2024.
- [84] Xichen Pan, Satya Narayan Shukla, Aashu Singh, Zhuokai Zhao, Shlok Kumar Mishra, Jialiang Wang, Zhiyang Xu, Jiahai Chen, Kunpeng Li, Felix Juefei-Xu, Ji Hou, and Saining Xie. Transfer between modalities with metaqueries. *arXiv preprint arXiv:2504.06256*, 2025.
- [85] Ziqi Pang, Tianyuan Zhang, Fujun Luan, Yunze Man, Hao Tan, Kai Zhang, William T. Freeman, and Yu-Xiong Wang. Randar: Decoder-only autoregressive visual generation in random orders. *arXiv preprint arXiv:2412.01827*, 2024.
- [86] William Peebles and Saining Xie. Scalable diffusion models with transformers. *arXiv preprint arXiv:2212.09748*, 2022.
- [87] William Peebles and Saining Xie. Scalable diffusion models with transformers. In *ICCV*, pages 4195–4205, 2023.

- [88] Guilherme Penedo, Quentin Malartic, Daniel Hesslow, Ruxandra Cojocaru, Hamza Alobeidli, Alessandro Cappelli, Baptiste Pannier, Ebtesam Almazrouei, and Julien Launay. The refinedweb dataset for falcon LLM: outperforming curated corpora with web data only. In *NeurIPS*, 2023.
- [89] Qi Qin, Le Zhuo, Yi Xin, Ruoyi Du, Zhen Li, Bin Fu, Yiting Lu, Xinyue Li, Dongyang Liu, Xiangyang Zhu, Will Beddow, Erwann Millon, Wenhai Wang Victor Perez, Yu Qiao, Bo Zhang, Xiaohong Liu, Hongsheng Li, Chang Xu, and Peng Gao. Lumina-image 2.0: A unified and efficient image generative framework, 2025.
- [90] Liao Qu, Huichao Zhang, Yiheng Liu, Xu Wang, Yi Jiang, Yiming Gao, Hu Ye, Daniel K Du, Zehuan Yuan, and Xinglong Wu. Tokenflow: Unified image tokenizer for multimodal understanding and generation. *arXiv preprint arXiv:2412.03069*, 2024.
- [91] Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, Gretchen Krueger, and Ilya Sutskever. Learning transferable visual models from natural language supervision. In *ICML*, pages 8748–8763, 2021.
- [92] Weiming Ren, Harry Yang, Ge Zhang, Cong Wei, Xinrun Du, Stephen Huang, and Wenhui Chen. Consist2v: Enhancing visual consistency for image-to-video generation. *arXiv preprint arXiv:2402.04324*, 2024.
- [93] Robin Rombach, Andreas Blattmann, Dominik Lorenz, Patrick Esser, and Björn Ommer. High-resolution image synthesis with latent diffusion models. In *CVPR*, pages 10684–10695, 2022.
- [94] Team Seaweed, Ceyuan Yang, Zhijie Lin, Yang Zhao, Shanchuan Lin, Zhibei Ma, Haoyuan Guo, Hao Chen, Lu Qi, Sen Wang, et al. Seaweed-7b: Cost-effective training of video generation foundation model. *arXiv preprint arXiv:2504.08685*, 2025.
- [95] Qingyu Shi, Jinbin Bai, Zhuoran Zhao, Wenhao Chai, Kaidong Yu, Jianzong Wu, Shuangyong Song, Yunhai Tong, Xiangtai Li, Xuelong Li, et al. Muddit: Liberating generation beyond text-to-image with a unified discrete diffusion model. *arXiv preprint arXiv:2505.23606*, 2025.
- [96] Weijia Shi, Xiaochuang Han, Chunting Zhou, Weixin Liang, Xi Victoria Lin, Luke Zettlemoyer, and Lili Yu. Lmfusion: Adapting pretrained language models for multimodal generation. *arXiv preprint arXiv:2412.15188*, 2024.
- [97] Wei Song, Yuran Wang, Zijia Song, Yadong Li, Haoze Sun, Weipeng Chen, Zenan Zhou, Jianhua Xu, Jiaqi Wang, and Kaicheng Yu. Dualtoken: Towards unifying visual understanding and generation with dual visual vocabularies. *arXiv preprint arXiv:2503.14324*, 2025.
- [98] Peize Sun, Yi Jiang, Shoufa Chen, Shilong Zhang, Bingyue Peng, Ping Luo, and Zehuan Yuan. Autoregressive model beats diffusion: Llama for scalable image generation. *arXiv preprint arXiv:2406.06525*, 2024.
- [99] Alexander Swerdlow, Mihir Prabhudesai, Siddharth Gandhi, Deepak Pathak, and Katerina Fragkiadaki. Unified multimodal discrete diffusion. *arXiv preprint arXiv:2503.20853*, 2025.
- [100] Hongxuan Tang, Hao Liu, and Xinyan Xiao. Ugen: Unified autoregressive multimodal model with progressive vocabulary learning. *arXiv preprint arXiv:2503.21193*, 2025.
- [101] Zineng Tang, Ziyi Yang, Mahmoud Khademi, Yang Liu, Chenguang Zhu, and Mohit Bansal. Codi-2: In-context, interleaved, and interactive any-to-any generation. 2023.
- [102] Zineng Tang, Ziyi Yang, Chenguang Zhu, Michael Zeng, and Mohit Bansal. Any-to-any generation via composable diffusion. *NeurIPS*, 36, 2024.
- [103] Chameleon Team. Chameleon: Mixed-modal early-fusion foundation models. *arXiv preprint arXiv:2405.09818*, 2024.
- [104] Gemini Team, Rohan Anil, Sebastian Borgeaud, Yonghui Wu, Jean-Baptiste Alayrac, Jiahui Yu, Radu Soricut, Johan Schalkwyk, Andrew M Dai, Anja Hauth, et al. Gemini: a family of highly capable multimodal models. *arXiv preprint arXiv:2312.11805*, 2023.
- [105] Rui Tian, Mingfei Gao, Mingze Xu, Jiaming Hu, Jiasen Lu, Zuxuan Wu, Yinfei Yang, and Afshin Dehghan. Unigen: Enhanced training & test-time strategies for unified multimodal understanding and generation. *arXiv preprint arXiv:2505.14682*, 2025.
- [106] Shengbang Tong, Ellis Brown, Penghao Wu, Sanghyun Woo, Manoj Middepogu, Sai Charitha Akula, Jihan Yang, Shusheng Yang, Adithya Iyer, Xichen Pan, Austin Wang, Rob Fergus, Yann LeCun, and Saining Xie. Cambrian-1: A fully open, vision-centric exploration of multimodal llms, 2024.
- [107] 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.
- [108] Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, Aurélien Rodriguez, Armand Joulin, Edouard Grave, and Guillaume Lample. Llama: Open and efficient foundation language models. *CoRR*, abs/2302.13971, 2023.
- [109] Ang Wang, Baole Ai, Bin Wen, Chaojie Mao, Chen-Wei Xie, Di Chen, Fei Wu, Haiming Zhao, Jianxiao Yang, Jianyuan Zeng, Jiayu Wang, Jingfeng Zhang, Jingren Zhou, Jinkai Wang, Jixuan Chen, Kai Zhu, Kang Zhao, Keyu Yan, Lianghua Huang, Mengyang Feng, Ningyi Zhang, Pandeng Li, Pingyu Wu, Ruihang Chu, Ruili Feng, Shiwei Zhang, Siyang Sun, Tao Fang, Tianxing Wang, Tianyi Gui,

- Tingyu Weng, Tong Shen, Wei Lin, Wei Wang, Wei Wang, Wenmeng Zhou, Wenten Wang, Wenting Shen, Wenyan Yu, Xianzhong Shi, Xiaoming Huang, Xin Xu, Yan Kou, Yangyu Lv, Yifei Li, Yijing Liu, Yiming Wang, Yingya Zhang, Yitong Huang, Yong Li, You Wu, Yu Liu, Yulin Pan, Yun Zheng, Yuntao Hong, Yupeng Shi, Yutong Feng, Zeyinzi Jiang, Zhen Han, Zhi-Fan Wu, and Ziyu Liu. Wan: Open and advanced large-scale video generative models. *arXiv preprint arXiv:2503.20314*, 2025.
- [110] Alex Jinpeng Wang, Dongxing Mao, Jiawei Zhang, Weiming Han, Zhuobai Dong, Linjie Li, Yiqi Lin, Zhengyuan Yang, Libo Qin, Fuwei Zhang, et al. Textatlas5m: A large-scale dataset for dense text image generation. *arXiv preprint arXiv:2502.07870*, 2025.
- [111] Bohan Wang, Zhongqi Yue, Fengda Zhang, Shuo Chen, Li'an Bi, Junzhe Zhang, Xue Song, Kennard Yanting Chan, Jiachun Pan, Weijia Wu, Mingze Zhou, Wang Lin, Kaihang Pan, Saining Zhang, Liyu Jia, Wentao Hu, Wei Zhao, and Hanwang Zhang. Discrete visual tokens of autoregression, by diffusion, and for reasoning. 2025.
- [112] Chunwei Wang, Guansong Lu, Junwei Yang, Runhui Huang, Jianhua Han, Lu Hou, Wei Zhang, and Hang Xu. ILLUME: illuminating your llms to see, draw, and self-enhance. *arXiv preprint arXiv:2412.06673*, 2024.
- [113] Jin Wang, Yao Lai, Aoxue Li, Shifeng Zhang, Jiacheng Sun, Ning Kang, Chengyue Wu, Zhenguo Li, and Ping Luo. Fudoki: Discrete flow-based unified understanding and generation via kinetic-optimal velocities. 2025.
- [114] Jiuniu Wang, Hangjie Yuan, Dayou Chen, Yingya Zhang, Xiang Wang, and Shiwei Zhang. Modelscape text-to-video technical report. *arXiv preprint arXiv:2308.06571*, 2023.
- [115] Xinlong Wang, Xiaosong Zhang, Zhengxiong Luo, Quan Sun, Yufeng Cui, Jinsheng Wang, Fan Zhang, Yuezhe Wang, Zhen Li, Qiyang Yu, et al. Emu3: Next-token prediction is all you need. *arXiv preprint arXiv:2409.18869*, 2024.
- [116] Yaohui Wang, Xinyuan Chen, Xin Ma, Shangchen Zhou, Ziqi Huang, Yi Wang, Ceyuan Yang, Yinan He, Jiashuo Yu, Peiqing Yang, et al. Lavie: High-quality video generation with cascaded latent diffusion models. *IJCV*, 2024.
- [117] Zekun Wang, King Zhu, Chunpu Xu, Wangchunshu Zhou, Jiaheng Liu, Yibo Zhang, Jiashuo Wang, Ning Shi, Siyu Li, Yizhi Li, Haoran Que, Zhaoxiang Zhang, Yuanxing Zhang, Ge Zhang, Ke Xu, Jie Fu, and Wenhao Huang. Mio: A foundation model on multimodal tokens. *arXiv preprint arXiv: 2409.17692*, 2024.
- [118] Chenyuan Wu, Pengfei Zheng, Ruirao Yan, Shitao Xiao, Xin Luo, Yuezhe Wang, Wanli Li, Xiyan Jiang, Yexin Liu, Junjie Zhou, Ze Liu, Ziyi Xia, Chaofan Li, Haoqiang Deng, Jiahao Wang, Kun Luo, Bo Zhang, Defu Lian, Xinlong Wang, Zhongyuan Wang, Tiejun Huang, and Zheng Liu. Omnigen2: Exploration to advanced multimodal generation. *arXiv preprint arXiv:2506.18871*, 2025.
- [119] Junfeng Wu, Yi Jiang, Chuofan Ma, Yuliang Liu, Hengshuang Zhao, Zehuan Yuan, Song Bai, and Xiang Bai. Liquid: Language models are scalable multi-modal generators. *arXiv preprint arXiv:2412.04332*, 2024.
- [120] Jay Zhangjie Wu, Yixiao Ge, Xintao Wang, Stan Weixian Lei, Yuchao Gu, Yufei Shi, Wynne Hsu, Ying Shan, Xiaohu Qie, and Mike Zheng Shou. Tune-a-video: One-shot tuning of image diffusion models for text-to-video generation. In *ICCV*, 2023.
- [121] Shengqiong Wu, Hao Fei, Leigang Qu, Wei Ji, and Tat-Seng Chua. Next-gpt: Any-to-any multimodal llm. *arXiv preprint arXiv:2309.05519*, 2023.
- [122] Size Wu, Zhonghua Wu, Zerui Gong, Qingyi Tao, Sheng Jin, Qinyue Li, Wei Li, and Chen Change Loy. Openuni: A simple baseline for unified multimodal understanding and generation. 2025.
- [123] Size Wu, Wenwei Zhang, Lumin Xu, Sheng Jin, Zhonghua Wu, Qingyi Tao, Wentao Liu, Wei Li, and Chen Change Loy. Harmonizing visual representations for unified multimodal understanding and generation, 2025.
- [124] Yecheng Wu, Zhuoyang Zhang, Junyu Chen, Haotian Tang, Dacheng Li, Yunhao Fang, Ligeng Zhu, Enze Xie, Hongxu Yin, Li Yi, et al. Vila-u: a unified foundation model integrating visual understanding and generation. *arXiv preprint arXiv:2409.04429*, 2024.
- [125] Yicheng Xiao, Lin Song, Rui Yang, Cheng Cheng, Zunnan Xu, Zhaoyang Zhang, Yixiao Ge, Xiu Li, and Ying Shan. Haploomni: Unified single transformer for multimodal video understanding and generation. *arXiv preprint arXiv:2506.02975*, 2025.
- [126] Enze Xie, Junsong Chen, Yuyang Zhao, Jincheng Yu, Ligeng Zhu, Yujun Lin, Zhekai Zhang, Muyang Li, Junyu Chen, Han Cai, et al. Sana 1.5: Efficient scaling of training-time and inference-time compute in linear diffusion transformer, 2025.
- [127] Enze Xie, Junsong Chen, Yuyang Zhao, Jincheng Yu, Ligeng Zhu, Yujun Lin, Zhekai Zhang, Muyang Li, Junyu Chen, Han Cai, et al. Sana 1.5: Efficient scaling of training-time and inference-time compute in linear diffusion transformer, 2025.
- [128] Jinheng Xie, Yuexiang Li, Yawen Huang, Haozhe Liu, Wentian Zhang, Yefeng Zheng, and Mike Zheng Shou. Boxdiff: Text-to-image synthesis with training-free box-constrained diffusion. In *ICCV*, pages 7452–7461, 2023.

- [129] Jinheng Xie, Weijia Mao, Zechen Bai, David Junhao Zhang, Weihao Wang, Kevin Qinghong Lin, Yuchao Gu, Zhijie Chen, Zhenheng Yang, and Mike Zheng Shou. Show-o: One single transformer to unify multimodal understanding and generation. In *ICLR*, 2025.
- [130] Rongchang Xie, Chen Du, Ping Song, and Chang Liu. MUSE-VL: modeling unified VLM through semantic discrete encoding. *arXiv preprint arXiv:2411.17762*, 2024.
- [131] Jinbo Xing, Menghan Xia, Yong Zhang, Haoxin Chen, Xintao Wang, Tien-Tsin Wong, and Ying Shan. Dynamicrafter: Animating open-domain images with video diffusion priors. *arXiv preprint arXiv:2310.12190*, 2023.
- [132] Lin Xu, Yilin Zhao, Daquan Zhou, Zhijie Lin, See Kiong Ng, and Jiashi Feng. Pllava: Parameter-free llava extension from images to videos for video dense captioning. *arXiv preprint arXiv:2404.16994*, 2024.
- [133] Zhiyang Xu, Jiuhai Chen, Zhaojiang Lin, Xichen Pan, Lifu Huang, Tianyi Zhou, Madian Khabza, Qifan Wang, Di Jin, Michihiro Yasunaga, et al. Pisces: An auto-regressive foundation model for image understanding and generation. *arXiv preprint arXiv:2506.10395*, 2025.
- [134] An Yang, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chengyuan Li, Dayiheng Liu, Fei Huang, Haoran Wei, Huan Lin, Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin Yang, Jiaxi Yang, Jingren Zhou, Junyang Lin, Kai Dang, Keming Lu, Keqin Bao, Kexin Yang, Le Yu, Mei Li, Mingfeng Xue, Pei Zhang, Qin Zhu, Rui Men, Runji Lin, Tianhao Li, Tingyu Xia, Xingzhang Ren, Xuancheng Ren, Yang Fan, Yang Su, Yichang Zhang, Yu Wan, Yuqiong Liu, Zeyu Cui, Zhenru Zhang, and Zihan Qiu. Qwen2.5 technical report. *arXiv preprint arXiv:2412.15115*, 2024.
- [135] Jian Yang, Dacheng Yin, Yizhou Zhou, Fengyun Rao, Wei Zhai, Yang Cao, and Zheng-Jun Zha. MMAR: towards lossless multi-modal auto-regressive probabilistic modeling. *arXiv preprint arXiv:2410.10798*, 2024.
- [136] Ling Yang, Ye Tian, Bowen Li, Xincheng Zhang, Ke Shen, Yunhai Tong, and Mengdi Wang. Mmada: Multimodal large diffusion language models. *arXiv preprint arXiv:2505.15809*, 2025.
- [137] Zhuoyi Yang, Jiayan Teng, Wendi Zheng, Ming Ding, Shiyu Huang, Jiazheng Xu, Yuanming Yang, Wenyi Hong, Xiaohan Zhang, Guanyu Feng, et al. Cogvideox: Text-to-video diffusion models with an expert transformer. *arXiv preprint arXiv:2408.06072*, 2024.
- [138] Xiang Yue, Yuansheng Ni, Tianyu Zheng, Kai Zhang, Ruoqi Liu, Ge Zhang, Samuel Stevens, Dongfu Jiang, Weiming Ren, Yuxuan Sun, Cong Wei, Botao Yu, Ruibin Yuan, Renliang Sun, Ming Yin, Boyuan Zheng, Zhenzhu Yang, Yibo Liu, Wenhao Huang, Huan Sun, Yu Su, and Wenhua Chen. MMMU: A massive multi-discipline multimodal understanding and reasoning benchmark for expert AGI. In *CVPR*, pages 9556–9567. IEEE, 2024.
- [139] Xiaohua Zhai, Basil Mustafa, Alexander Kolesnikov, and Lucas Beyer. Sigmoid loss for language image pre-training, 2023.
- [140] Biao Zhang and Rico Sennrich. Root mean square layer normalization. In *NeurIPS*, 2019.
- [141] David Junhao Zhang, Jay Zhangjie Wu, Jia-Wei Liu, Rui Zhao, Lingmin Ran, Yuchao Gu, Difei Gao, and Mike Zheng Shou. Show-1: Marrying pixel and latent diffusion models for text-to-video generation. *arXiv preprint arXiv:2309.15818*, 2023.
- [142] Hong Zhang, Zhongjie Duan, Xingjun Wang, Yingda Chen, Yuze Zhao, and Yu Zhang. Nexus-gen: A unified model for image understanding, generation, and editing. *arXiv preprint arXiv:2504.21356*, 2025.
- [143] Pan Zhang, Xiaoyi Dong, Yuhang Zang, Yuhang Cao, Rui Qian, Lin Chen, Qipeng Guo, Haodong Duan, Bin Wang, Linke Ouyang, et al. Internlm-xcomposer-2.5: A versatile large vision language model supporting long-contextual input and output. *arXiv preprint arXiv:2407.03320*, 2024.
- [144] Peiyuan Zhang, Kaichen Zhang, Bo Li, Guangtao Zeng, Jingkan Yang, Yuanhan Zhang, Ziyue Wang, Haoran Tan, Chunyuan Li, and Ziwei Liu. Long context transfer from language to vision. *arXiv preprint arXiv:2406.16852*, 2024.
- [145] Shiwei Zhang, Jiayu Wang, Yingya Zhang, Kang Zhao, Hangjie Yuan, Zhiwu Qin, Xiang Wang, Deli Zhao, and Jingren Zhou. I2vgen-xl: High-quality image-to-video synthesis via cascaded diffusion models. *arXiv preprint arXiv:2311.04145*, 2023.
- [146] Yuanhan Zhang, Jinming Wu, Wei Li, Bo Li, Zejun Ma, Ziwei Liu, and Chunyuan Li. Video instruction tuning with synthetic data. *arXiv preprint arXiv:2410.02713*, 2024.
- [147] Chuyang Zhao, Yuxing Song, Wenhao Wang, Haocheng Feng, Errui Ding, Yifan Sun, Xinyan Xiao, and Jingdong Wang. Monoformer: One transformer for both diffusion and autoregression. *arXiv preprint arXiv:2409.16280*, 2024.
- [148] Chunting Zhou, LILI YU, Arun Babu, Kushal Tirumala, Michihiro Yasunaga, Leonid Shamir, Jacob Kahn, Xuezhe Ma, Luke Zettlemoyer, and Omer Levy. Transfusion: Predict the next token and diffuse images with one multi-modal model. In *ICLR*, 2025.
- [149] Xianwei Zhuang, Yuxin Xie, Yufan Deng, Liming Liang, Jinghan Ru, Yuguo Yin, and Yuexian Zou. Vargpt: Unified understanding and generation in a visual autoregressive multimodal large language model, 2025.
- [150] Jialv Zou, Bencheng Liao, Qian Zhang, Wenyu Liu, and Xinggang Wang. Omnimamba: Efficient and unified multimodal understanding and generation via state space models, 2025.

NeurIPS Paper Checklist

The checklist is designed to encourage best practices for responsible machine learning research, addressing issues of reproducibility, transparency, research ethics, and societal impact. Do not remove the checklist: **The papers not including the checklist will be desk rejected.** The checklist should follow the references and follow the (optional) supplemental material. The checklist does NOT count towards the page limit.

Please read the checklist guidelines carefully for information on how to answer these questions. For each question in the checklist:

- You should answer [Yes], [No], or [NA].
- [NA] means either that the question is Not Applicable for that particular paper or the relevant information is Not Available.
- Please provide a short (1–2 sentence) justification right after your answer (even for NA).

The checklist answers are an integral part of your paper submission. They are visible to the reviewers, area chairs, senior area chairs, and ethics reviewers. You will be asked to also include it (after eventual revisions) with the final version of your paper, and its final version will be published with the paper.

The reviewers of your paper will be asked to use the checklist as one of the factors in their evaluation. While "[Yes]" is generally preferable to "[No]", it is perfectly acceptable to answer "[No]" provided a proper justification is given (e.g., "error bars are not reported because it would be too computationally expensive" or "we were unable to find the license for the dataset we used"). In general, answering "[No]" or "[NA]" is not grounds for rejection. While the questions are phrased in a binary way, we acknowledge that the true answer is often more nuanced, so please just use your best judgment and write a justification to elaborate. All supporting evidence can appear either in the main paper or the supplemental material, provided in appendix. If you answer [Yes] to a question, in the justification please point to the section(s) where related material for the question can be found.

IMPORTANT, please:

- **Delete this instruction block, but keep the section heading “NeurIPS Paper Checklist”,**
- **Keep the checklist subsection headings, questions/answers and guidelines below.**
- **Do not modify the questions and only use the provided macros for your answers.**

1. Claims

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

Answer: [Yes]

Justification: Claims in abstract and Section 1 accurately reflect the paper’s contributions and scope.

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 discuss the limitations of our work in Section A.4.

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 was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [NA]

Justification: There are no theoretical results in the paper.

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: In the experiment section, we include implementation details to ensure reproducibility. Besides, we will release the training code.

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
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

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: [Yes]

Justification: We release the training and inference code at <https://github.com/showlab/Show-o> and most of the training data is publicly available.

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, hyper-parameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

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

Justification: We provide implementation details in Section 4.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

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

Justification: As described in our Ethics Statement, we have conducted repeated evaluations and observed stable results with minimal variations. However, our experimental comparisons follow prior literature that did not report error bars. For language-only evaluations in the ablation studies, we explicitly report the standard deviations.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- 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 provide experiments compute resources in Section 4.

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]

Justification: We have carefully reviewed the Code of Ethics and strictly adhered to it in our work.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: We discuss the broader impacts in Section A.4.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- 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: [Yes]

Justification: For the dataset, we only provide the publicly available links to them. We add a safety checker to our inference pipeline to mitigate potential 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.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: The code, data, models are all properly credited and their license and terms of use are properly respected in our work.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- 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: [No]

Justification: We do not introduce new assets.

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: We do not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional review board (IRB) approvals or equivalent for research with human subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: We do not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- 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 rigor, or originality of the research, declaration is not required.

Answer: [Yes]

Justification: Our models is based on open-sourced LLM models, and we have claimed it in the methodology and experiments sections.

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.

A Technical Appendices and Supplementary Material

A.1 Experimental Setup

Datasets. The curated approximately 66M image-text pairs consist of images with a resolution of at least 512 pixels in width and height. The images are filtered from CC12M [14], COYO [12], LAION-Aesthetic-12M* and AI synthetic data. The images are recaptioned by LMMs except for the synthetic data. The 9M high-quality multimodal understanding instruction data is curated from Densefusion-1M [63], and LLaVA-OneVision [56].

Implementation Details. The semantic layers $\mathcal{S}(\cdot)$ are pre-distilled from SigLIP-so400m-patch14-384* over 200K iterations, using a batch size of 512 and a cosine-scheduled learning rate of $2e-5$. During distillation, Eq. 1 is applied to the visual latents with only a probability of 0.3 in the last 20K iterations. The input image resolution of 3D causal VAE encoder with 2×2 patch embedding layer is set as 432×432 to get $729 = 27 \times 27$ visual latents, which matches the ones extracted by SigLIP. Once distilled, the semantic layers $\mathcal{S}(\cdot)$ are capable of extracting rich semantic features from both clean and noised visual latents. In statistics, the extracted features from clean visual latents by $\mathcal{S}(\cdot)$ have converged to an average cosine similarity of around 0.9 with those extracted by the original SigLIP on the curated 66M image-text pairs. We interpolate the position embeddings in the bicubic mode when involving other image/video resolutions.

Our models build upon two LLM variants, *i.e.*, Qwen2.5-1.5B-Instruct [134] and Qwen2.5-7B-Instruct [134], respectively. We adopt 3D causal VAE proposed in Wan2.1 [109] with $8 \times$ and $4 \times$ spatial and temporal compression, respectively. In stage 1, we first train the 1.5B variant for 150K iterations using AdamW optimizer with a constant learning rate of 0.0001 on the curated 66M image-text pairs in a resolution of 432×432 . The context length of single image-text pairs is set as 1024. The total batch sizes for multimodal understanding and generation are 128 and 384, respectively. α in Eq. 4 is set as 0.2. For visual generation data, the caption is dropped with a probability of 0.1 to enable the classifier-free guidance. This training process roughly takes one and a half days using 64 H100 GPUs. Subsequently, we replace the generation data with 16M high-quality data (filtered from 66M image-text pairs) and continue to train for 40K iterations. In stage 2, we follow the training strategies in LLaVA-OneVision [56] to train the 1.5B model using around 9M multimodal instructional and 16M high-quality generation data for a total of around 35K iterations. α in Eq. 4 is set as 1.0. The stage 2 training process takes around 15 hours. For models with mixed-modality and video generation capabilities, we progressively add video-text and interleaved data in stage 1. For video data, we randomly sample a 2s 480p or 432×432 clips with 17 frames from each video with an interval of 3 frames. The context length at this time is set as 7006. In stage 2, high-quality video-text and interleaved data are added to further improve video and mixed-modality generation capabilities.

To further improve the image generation and text rendering quality, we further train the small-scale model on images with higher resolution (512×512 and 1024×1024) and involve an additional text-rich image data, *i.e.*, a subset of TextAtlas [110].

Building on the pre-trained image-level Show-o2 models, we enhance their video understanding capabilities by further fine-tuning on 1.6M video samples from [146], together with 1.1M image-level samples from the earlier stage. We adopt the same video training and inference settings as LLaVA-OneVision. The evaluation results are shown in Table 4.

In the training of our model based on the 7B LLM variant, we resume the flow head pre-trained based on the 1.5B model and additionally train the newly initialized spatial (-temporal) fusion, projector, and MLP transformations for 3K iterations with 2K warm-up steps to align the hidden size and then further train spatial (-temporal) fusion, the projector, MLP transformations, and the flow head together. Following that, we conduct the training stages 1 and 2 in the same manner as those of the 1.5B model. The whole training process of our 7B model takes approximately 2 and a half days using 128 H100 GPUs. We do not include interleaved and video data in the training stages of the larger model due to the huge computational cost and training duration.

* <https://huggingface.co/datasets/dclure/laion-aesthetics-12m-umap>

* <https://huggingface.co/google/siglip-so400m-patch14-384>

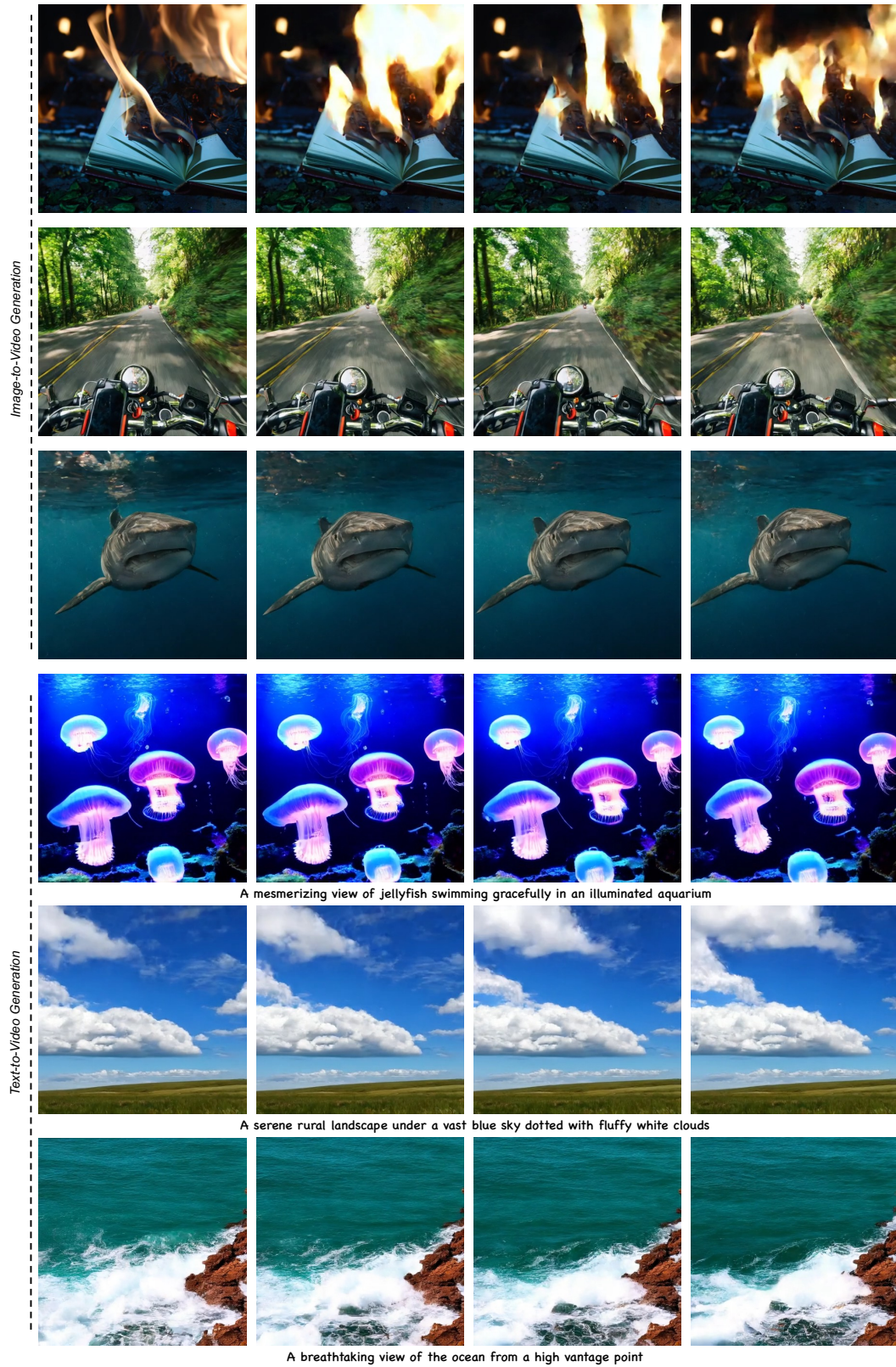


Figure 3: Text-to-video and image-to-video generation examples.



What is the unusual aspect in this video?

The unusual aspect in this video is the formation of a large amount of white smoke from a small paper, creating a dramatic and visually captivating cloud.



What's funny about this video?

LLaVA-OV-7B:

The video is funny because the person in the white shirt and gray shorts is seen dancing energetically, while the person in the blue shirt remains mostly motionless, creating a humorous contrast.

Show-o2-7B (Ours):

The funny part is that the person in the white shirt is singing or speaking into the microphone, while everyone else is sitting on the couch and looking at the person, possibly in amusement or confusion.



What's funny about this video?

LLaVA-OV-7B:

The funny part is the dog's reaction to the snake, which it tries to catch and then playfully interacts with.

Show-o2-7B (Ours):

The dog's exaggerated reaction to the snakes makes the scene comical.

Figure 4: Video understanding examples.

A.2 More Qualitative Results

Figures 3 and 4 provide more visualization results of video generation and understanding.

A.3 Text Prompts

We provide the text prompts for image generation used in Fig. 2 below:

“Hyper-detailed image of a mature man with short, graying hair and deep blue eyes. He has a rugged, weathered face with a strong jawline and a slight beard. His expression is thoughtful and introspective. The lighting is dramatic, highlighting the contours of his face. The photo is in 8K resolution, capturing every wrinkle and pore. ”

“A soft, natural portrait photograph captures a young woman with fair skin and long, ash-blonde hair cascading gently over her shoulders, her striking light blue eyes subtly enhanced with natural makeup and a gentle, calm smile playing on her lips. She wears a cozy, cream-colored winter sweater and a delicate woolen scarf adorned with subtle snowflake patterns, positioned slightly off-center, creating a sense of relaxed elegance. Behind her, a softly blurred snowy Moscow street scene unfolds, with traditional architecture and the diffused, golden glow of a winter afternoon contributing to a serene and contemplative atmosphere. At the very bottom of the frame, in simple, elegant lettering, appears the phrase "BE KIND". ”

“A vibrant, highly detailed close-up of a colorful parrot perched on a branch, featuring intricate feather textures, vivid colors (red, blue, green, yellow), and a tropical rainforest background. The parrot’s

Table 15: Comparative analysis of selected unified multimodal models based on the utilization of visual representations and type of unified modeling for multimodal understanding and generation. In this context, **native und. & gen.** refers to the direct decoding of output sequences into texts, images, and videos, as opposed to serving as conditions for decoding using external pre-trained decoders like Stable Diffusion. * indicates the method uses two distinct models for multimodal understanding and generation, respectively.

Methods	Und. & Gen. Representation			Type of Unified Modeling		
	Unified	Decoupled	Support Video	Native Und. & Gen.	Assembling Tailored Models	Paradigm
Chameleon [103]	✓		✗	✓		AR
Show-o [129]	✓		✗	✓		AR + Diff.
Transfusion [148]	✓		✗	✓		AR + Diff.
VILA-U [124]	✓		✓	✓		AR
Emu3 [115]	✓		✓	✓		AR
MonoFormer [147]	✓		✗	✓		AR + Diff.
Dual-Diffusion [64]	✓		✗	✓		Diff.
SynerGen-VL [58]	✓		✗	✓		AR
MMAR [135]	✓		✗	✓		AR + MAR
MUSE-VL [130]	✓		✗	✓		AR
Orthus [53]	✓		✗	✓		AR + Diff.
Liquid [119]	✓		✗	✓		AR
LlamaFusion [96]	✓		✗	✓		AR + Diff.
UGen [100]	✓		✗	✓		AR
UniDisc [99]	✓		✗	✓		Diff.
UniToken [50]	✓		✗	✓		AR
Harmon [123]	✓		✗	✓		AR+MAR
DualToken [97]	✓		✗	✓		AR
UniTok [78]	✓		✗	✓		AR
Selftok [111]	✓		✗	✓		AR
Muddit [95]	✓		✗	✓		Diff.
MMaDA [136]	✓		✗	✓		Diff.
HaploOmni [125]	✓		✓	✓		AR + Diff.
TokLIP [69]	✓		✗	✓		AR
Show-o2 (Ours)	✓		✓	✓		AR + Diff.
Janus-Series [26, 27, 80]		✓	✗	✓		AR (+Diff.)
VARGPT [149]		✓	✗	✓		AR
UnidFluid [38]		✓	✗	✓		AR + MAR
OmniMamba [150]		✓	✗	✓		AR
Mogao [66]		✓	✗	✓		AR + Diff.
BAGEL [32]		✓	✓	✓		AR + Diff.
Fudoki [113]		✓	✗	✓		Diff.
UniGen [105]		✓	✗	✓		AR + Diff.
NEXT-GPT [121]		✓	✓		✓	AR + Diff.
CoDI [102]		✓	✓		✓	AR + Diff.
DreamLLM [36]		✓	✗		✓	AR + Diff.
SEED-X [40]		✓	✗		✓	AR + Diff.
MIO [117]		✓	✓		✓	AR + Diff.
CoDI-2 [101]		✓	✓		✓	AR + Diff.
MetaMorph [107]		✓	✗		✓	AR + Diff.
ILLUME [112]		✓	✗		✓	AR + Diff.
ILLUME+ [46]		✓	✗		✓	AR + Diff.
MetaQueries [84]		✓	✗		✓	AR + Diff.
Nexus-Gen [142]		✓	✗		✓	AR + Diff.
Ming-Lite-Uni [42]		✓	✗		✓	AR + Diff.
BLIP3-o [18]		✓	✗		✓	AR + Diff.
OpenUni [122]		✓	✗		✓	AR + Diff.
UniWorld [68]		✓	✗		✓	AR + Diff.
Ming-Omni [5]		✓	✓		✓	AR + Diff.
Pisces [133]		✓	✗		✓	AR + Diff.
TokenFlow* [90]	✓		✗		✓	AR
SemHiTok* [28]	✓		✗		✓	AR

eyes are sharp and expressive, with a natural glint of light. The image is photorealistic, ultra HD (8K resolution), with soft natural lighting and a shallow depth of field, creating a blurred bokeh effect in the background. The scene is peaceful and lush, showcasing the beauty of nature. ”

“A dark, moody room with a glowing neon sign on the wall that spells out ’SHOW O2’ in bold, vibrant pink and blue colors. The neon light reflects softly on the polished concrete floor, creating a futuristic and artistic vibe. ”

A.4 Limitations and Broader Impacts

We found that our model is not good at rendering text on the image. We investigated our generation datasets and observed that the proportion of images with rendered texts is relatively small, which

potentially leads to bad text rendering. In addition, the generated images will lack details of the small objects because of the limited image resolution. To address this limitation, as outlined in the implementation details, we have enhanced the model by training it on higher-resolution data and incorporating image datasets rich in textual information.

Our models possess the ability to generate text and images, which may carry the risk of unintended misuse. As discussed in our Ethics Statement, we have conducted data filtering to avoid this issue and integrated a safety checker into our inference pipeline to mitigate potential misuse.