

LlamaSeg: Image Segmentation via Autoregressive Mask Generation

Jiru Deng^{1,2*}, Tengjin Weng^{1,3*}, Tianyu Yang⁴, Wenhan Luo⁵, Zhiheng Li²,
and Wenhao Jiang^{1†}

¹ Guangdong Laboratory of Artificial Intelligence and Digital Economy (SZ)

² Tsinghua University

³ Shenzhen University

⁴ Meituan

⁵ Division of AMC and Department of ECE, HKUST

djr23@mails.tsinghua.edu.cn, wtjdsb@gmail.com, tianyu-yang@outlook.com,
whluo.china@gmail.com, zhhli@mail.tsinghua.edu.cn, cswjiang@gmail.com

Abstract. We present **LlamaSeg**, a visual autoregressive framework that unifies multiple image segmentation tasks via natural language instructions. By reformulating segmentation as visual generation, LlamaSeg encodes masks as visual tokens and uses a LLaMA-style Transformer for direct next-token prediction, naturally fitting segmentation into autoregressive architectures. To support large-scale training, we introduce a data annotation pipeline and construct the **SA-OVRS** dataset, which contains **2M** segmentation masks annotated with over **5,800** open vocabulary labels or diverse textual descriptions, spanning diverse real-world scenarios. This enables our model to localize objects in images based on text prompts and to generate fine-grained masks. We further introduce the composite metric average Hausdorff Distance (d_{AHD}) to evaluate mask contour fidelity for generative models better. Experiments show that LlamaSeg consistently outperforms existing generative approaches on multiple segmentation benchmarks and delivers finer, more accurate segmentation masks. Code and dataset are available at <https://github.com/GML-FMGroup/llamaseg>.

1 Introduction

Recent advances in multimodal large language models (MLLMs) have demonstrated the potential of unifying diverse vision-language tasks under an autoregressive framework [1, 9, 55, 61]. However, extending this paradigm to dense segmentation remains challenging, as it demands precise pixel-level generation [46, 52]. As illustrated in Figure 1, existing autoregressive segmentation approaches can be broadly grouped into three paradigms. Embedding-based representations [18, 53] predict a latent mask token such as <SEG> and then delegate pixel-level mask generation to a specialist model like SAM [17]. Although effective, this

* Equal contribution.

† Corresponding author.

reliance on external segmentation experts introduces task-specific components and undermines the goal of a unified end-to-end framework. Coordinate-driven modeling [29, 46] encodes each segmentation mask as a sequence of polygon vertices. While naturally aligned with next-token prediction, it requires long coordinate sequences to capture intricate boundaries, reducing efficiency and scalability for dense semantic segmentation. Text-conditioned generation [19, 43] formulates segmentation as a text generation task by emitting category tokens or free-form descriptions that must later be converted into pixel masks, which inevitably limits spatial precision. These limitations highlight that current autoregressive methods remain inadequate for fine-grained, end-to-end segmentation within a single unified framework.

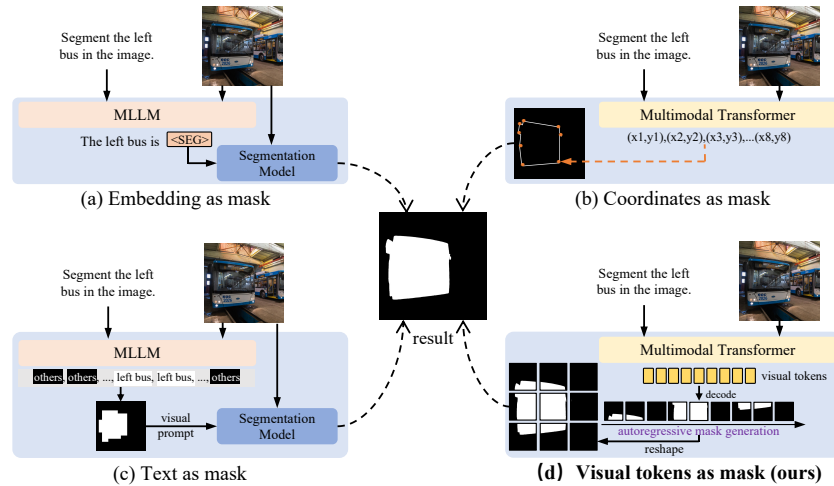


Fig. 1: Comparison of autoregressive segmentation paradigms. (a) **Embedding as mask.** (b) **Coordinates as mask.** (c) **Text as mask.** (d) **Visual tokens as mask (ours)**, which directly autoregresses visual tokens to capture fine-grained structures and yields precise pixel-level masks.

However, the progress of autoregressive image generation [12, 32, 35] shows that appropriate representations and objectives enable these frameworks to model complex visuals. These methods, which typically condition image synthesis on category labels, text prompts, or visual cues, highlight the expressive power of token-based modeling. Their potential remains underexplored in structured prediction tasks with deterministic ground truth, such as segmentation. **Building on these insights, we aim to bridge this gap by revisiting autoregressive modeling for segmentation and devising a representation that predicts high-fidelity masks as discrete visual tokens within a unified next-token prediction framework.**

In this work, we propose **LlamaSeg**, which reformulates image segmentation as a visual generation task, as shown in panel (d) of Figure 1. We treat segmentation masks as image-like structures, which are discretized into token sequences via VQGAN [12]. A LLaMA-based autoregressive model [38, 39] then generates mask tokens conditioned on the encoded visual inputs, maintaining the standard next-token prediction paradigm, thereby facilitating the integration of segmentation tasks into general autoregressive frameworks. Our framework is scalable across different model sizes and image resolutions, and supports training from scratch as well as fine-tuning on pretrained MLLMs. To address the lack of suitable training data and evaluation tools for this paradigm, we introduce a large-scale dataset, SA-OVRS, containing 2 million segmentation instances annotated with open-vocabulary labels or rich textual descriptions derived from SA-1B [17], and propose the average Hausdorff Distance (d_{AHD}) under multi-level IoU thresholds as a new metric to assess contour fidelity in generated masks. Experiments on multiple segmentation benchmarks demonstrate that our model produces high-quality, fine-grained masks and outperforms existing visual generative approaches. Our main contributions are as follows:

1. We propose **LlamaSeg**, a visual autoregressive segmentation framework that shares the core next-token generation paradigm of Large Language Models (LLMs), allowing segmentation to be natively unified within the LLMs architecture while effectively modeling visual context and producing high-quality, fine-grained masks.
2. We develop a data annotation pipeline and construct the **SA-OVRS** dataset comprising 2M masks with over 5,800 open-vocabulary or richly descriptive labels. The dataset and the associated annotation pipeline are released as reusable tools to advance future segmentation research and to facilitate the creation of new high-quality datasets.
3. We introduce a composite evaluation metric d_{AHD} to assess mask-contour fidelity and demonstrate that LlamaSeg surpasses most existing visual generative models on diverse semantic and referring segmentation benchmarks, delivering more precise pixel-level masks.

2 Related Work

2.1 Multimodal Large Language Models

The research on MLLMs is closely related to cross-modal representation learning. With the rapid advancement of LLMs [3, 8, 30, 42], recent efforts have increasingly focused on aligning visual representations with the language embedding space. BLIP-2 [21] and InstructBLIP [9] incorporate lightweight Querying Transformer modules to inject visual context. The LLaVA series [20, 23, 24, 50, 51] employs simple linear projection layers to map image features. The Qwen-VL family [2, 44] leverages cross-attention mechanisms with learnable queries for tighter vision-language integration. These approaches collectively enable large language models to interpret and reason over visual content effectively. Beyond

image-level understanding, recent studies have begun to extend vision-language models toward dense prediction tasks such as image segmentation, which require fine-grained spatial reasoning and localized visual grounding.

2.2 Image Segmentation

Transformer-based segmentation method. Segmentation models based on the Transformer [42] architecture typically include MaskFormer [7], SegFormer [54], and SAM [17]. Since language models such as BERT [10] also adopt Transformer architectures, this shared foundation facilitates cross-modal information fusion. Several approaches [11, 22, 48] have explored referring image segmentation built upon this architecture. Recent studies have increasingly focused on MLLMs. For example, LISA [18] integrates SAM and LLaVA [24], leveraging token embeddings to guide the segmentation model. Subsequent methods [33, 53] build upon this framework with further enhancements. Despite their strong performance, these approaches typically adopt a pipeline architecture, which increases model complexity and training difficulty.

Generative segmentation method. Generative methods recast segmentation as a mask-generation problem conditioned on the input image, providing a streamlined alternative to traditional pipeline-based approaches. GSS [5], for example, links the mask posterior distribution with the latent prior of the input image for semantic segmentation. Despite adopting the generative paradigm, these approaches remain confined to vision-only tasks. Unified-IO [28] and Unified-IO2 [27] extend to multiple modalities and include segmentation, but neither explores it deeply. Unified-IO struggles with complex language understanding, while Unified-IO2 employs a two-stage procedure dependent on bounding boxes. Such reliance on task-specific modules or intermediate representations hinders seamless end-to-end segmentation and limits generalization.

2.3 Autoregressive Image Generation

The field of image generation has witnessed significant advancements since autoregressive modeling achieved breakthroughs in NLP. Numerous high-quality models [31, 36, 37, 45, 47, 58] have emerged in recent years. These models employ image tokenizers [12, 32, 41] to convert images into discrete representations, which are then reshaped into 1D sequences to align with the autoregressive next-token prediction framework. Subsequent research [34, 35, 57, 62, 63] has focused on designing image tokenizers with improved generation quality, and on comparing them with continuous distribution models such as VAE [16] to demonstrate their competitiveness. In our work, we adopt VQGAN [12] to encode masks into discrete representations. An image tokenizer trained on a large corpus of colorful images is sufficient for mask modeling and is capable of preserving contour details.

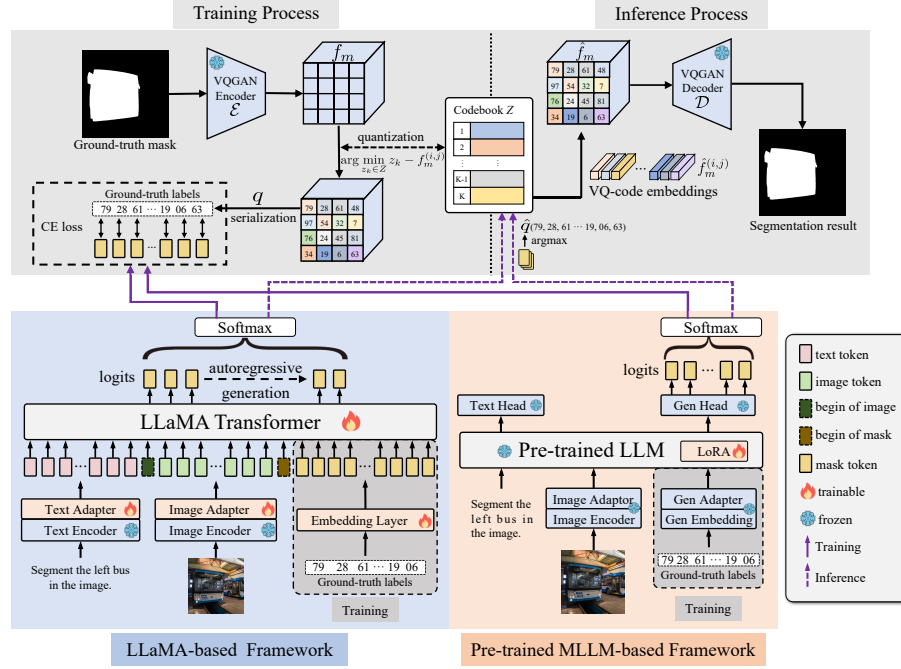


Fig. 2: Overall framework of **LlamaSeg**. The model consists of a VQGAN-based mask tokenizer and a LLaMA-style autoregressive generator. During training, ground-truth masks are quantized into discrete codes to supervise token prediction; during inference, predicted tokens are decoded to produce segmentation masks. We support training either from scratch or by adapting a pre-trained MLLM with lightweight modules.

3 LlamaSeg

3.1 Overview

The overall framework of LlamaSeg is illustrated in Fig. 2. For the LLaMA-based framework, we employ image and text encoders to extract input features, which are projected to align with the embedding dimension of LLaMA [38, 39] and concatenated with learnable separator tokens to form the input sequence. The LLaMA model then autoregressively generates mask tokens as output. We adopt VQGAN [12] as the mask tokenizer. During training, the mask tokenizer encodes the ground truth mask into code indices, which serve as the training targets for the LLaMA model. During inference, the mask tokenizer receives the mask tokens generated by the LLaMA model, retrieves the corresponding codes from the codebook, and decodes them to produce the segmentation mask. For the MLLM-based framework utilizing a pretrained MLLM, the text encoder is omitted, while the remaining procedure remains unchanged.

3.2 Mask Tokenizer

We use an image tokenizer to transform segmentation masks into discrete tokens. Specifically, we use a VQGAN from LlamaGen [35] with a downsample rate of 16 and a codebook $Z \in \mathbb{R}^{K \times d_{vq}}$ containing K vectors, where d_{vq} is the vector dimension. We consider the segmentation mask as a special RGB image composed only of black and white pixels. During training, the encoder $\mathcal{E}$ transforms the normalized mask $M_I \in \mathbb{R}^{H \times W \times 3}$ into features $f_m \in \mathbb{R}^{h \times w \times d_{vq}}$, and the quantizer $\mathcal{Q}$ converts features into discrete tokens $q \in [K]^{h \times w}$. The quantizer matches each vector in f_m with the closest vector in Z in terms of Euclidean distance and finds the corresponding code index:

$$q^{(i,j)} = \underset{z_k \in Z}{\text{argmin}} \|z_k - f_m^{(i,j)}\| \in [K], \quad (1)$$

where z_k is the vector within Z . Flattening the code indices to form a 1D sequence can serve as supervised training data for an autoregressive model. At inference time, for a 1D token output by an autoregressive model, the token IDs are equivalent to the code indices in the mask tokenizer. By “looking up in table”, vectors $\hat{z}_k \in \mathbb{R}^{d_{vq}}$ can be found in the codebook according to $\hat{q} \in [K]^{h \times w}$, which are rearranged into a 2D shape from token IDs and form the feature map $\hat{f}_m \in \mathbb{R}^{h \times w \times d_{vq}}$. The normalized mask image $\hat{M}_I \in [-1, 1] \subset \mathbb{R}^{H \times W \times 3}$ is reconstructed by the decoder $\mathcal{D}$ using $\hat{f}_m$. Average $\hat{M}_I$ across the channel dimension and then apply a binarization step to produce the final segmentation result $\hat{M}_{bin}^{(i,j)}$:

$$\hat{M} = \frac{1}{3} \sum_{i=0}^2 (\hat{M}_I[:, :, i]), \quad \hat{M}_{bin}^{(i,j)} = \begin{cases} 1, & \hat{M}^{(i,j)} \geq 0 \\ 0, & \text{otherwise} \end{cases}. \quad (2)$$

3.3 Image and Text Feature Extraction

We construct textual inputs by combining object labels or descriptions with 10 predefined templates, such as “Produce a segmentation mask for the $\{object\}$ name.” For the autoregressive model trained from scratch, we employ a frozen SigLIP2 [40] with a patch size of 16 to encode both images and text. The vision encoder’s patch size matches the mask tokenizer’s downsampling rate, ensuring that input tokens correspond to the same pixel regions as the segmentation mask for precise token-level alignment. Furthermore, trainable adaptors are used to project the image and text features into the embedding dimension of the autoregressive model. Each adaptor has two linear layers and a GELU [13] activation function. We consider visual encoders with input resolutions of 256 and 384, corresponding to 256 and 576 visual tokens, respectively.

3.4 Autoregressive Model

LLaMA Structure. The autoregressive model is based on LLaMA [38, 39] model. Each Transformer layer employs 1D RoPE, treating both images and

masks as 1D sequences. We utilize two model variants of different sizes. The base model comprises 770 million parameters, 16 layers, a hidden size of 1920, and 20 attention heads. The large model contains 1.5 billion parameters, 22 layers, a hidden size of 2304, and 32 attention heads. During training, input image and text features are concatenated along the channel dimension. To distinguish between text, image, and mask tokens, we introduce learnable separators: `<BOI>` (*Begin-Of-Image*) and `<BOM>` (*Begin-Of-Mask*). The input sequence during training is structured as follows:

`[text tokens]<BOI>[image tokens]<BOM>[mask tokens]`.

The model is trained using a cross-entropy loss, computed solely on the mask tokens. At inference time, only `<BOM>` and the preceding tokens are fed into the model, allowing it to generate mask tokens autoregressively.

Utilization of pre-trained MLLM. We adopt Janus Pro [6], an MLLM for multimodal understanding and generation with separate language and image pathways. For multimodal understanding, images are encoded by the SigLIP [60] vision encoder. For image generation, however, visual inputs are synthesized via an image tokenizer [12]. The model operates at an image resolution of 384. To enable visual generation, we utilize components responsible for producing mask tokens, denoted as “Gen Embedding”, “Gen Adapter”, and “Gen Head” in Fig. 2. These components are kept frozen to retain the generative capabilities acquired through large-scale data training. During training, the sequence input to the model is:

`<|User|>: [text tokens]<image> <|Assistant|>: <begin_of_mask>[mask tokens]`

where `<|User|>` and `<|Assistant|>` denote the dialogue roles, `<image>` represents the input image, and `<begin_of_mask>` indicates the beginning of the mask token sequence. During inference, only `<begin_of_mask>` and its preceding context are provided as input.

4 SA-OVRS Dataset Construction

Existing methods [18, 53] leverage semantic and referring segmentation datasets for language-guided segmentation. Semantic datasets such as COCO-Stuff [4] treat category names as text, but hierarchical labels (e.g., “vegetables” vs. “broccoli”) are split into independent classes, leading to semantic inconsistencies and limiting open-vocabulary learning. In addition, existing datasets are too small to support the dense cross-modal alignment required by autoregressive models. Figure 3 illustrates the two-stage annotation pipeline used to construct SA-OVRS. In the first stage, we generate candidate open-vocabulary labels for SA-1B [17] images and match them with their corresponding masks through quality filtering, while the second stage generates natural-language descriptions for each object instance, encompassing both referring and reasoning expressions. More details are presented in the following subsection.

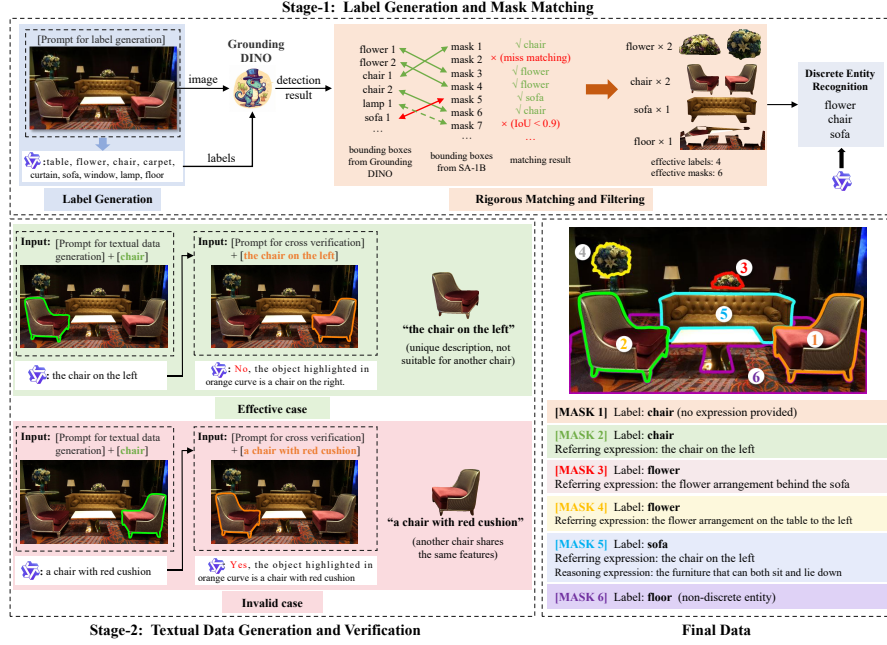


Fig. 3: A two-stage annotation pipeline for constructing the **SA-OVRS** dataset. **Stage 1: Label Generation and Mask Matching (top)** leverages Qwen2-VL to generate candidate open-vocabulary labels, which are then aligned with SA-1B masks using GroundingDINO through strict matching and filtering procedures. **Stage 2: Textual Data Generation and Verification (bottom left)** produces high-quality referring expressions and reasoning descriptions.

4.1 Label Generation and Mask Matching

For each SA-1B image, we first obtain candidate open-vocabulary labels using Qwen2-VL-72B with a tailored prompt, limiting outputs to at most ten labels per image. GroundingDINO [25] then detects bounding boxes for each label. To ensure high-quality matches, we apply a three-step filtering procedure: (i) Remove labels associated with excessive boxes (more than four). (ii) Eliminate nested detections where a smaller box is almost entirely enclosed in a larger one ($\text{IoU} > 0.97$). (iii) Keep only matches with sufficient confidence (score > 0.3) and strong mask overlap ($\text{IoU} > 0.85$ for single-box labels, 0.9 for multi-box labels). Masks linked to the same label are merged to produce semantic segmentation samples.

4.2 Textual Data Generation and Verification

To produce unique natural-language descriptions for each instance, we highlight every matched mask with a green contour and prompt Qwen2-VL-72B to gen-

erate a referring expression that distinguishes it from all other objects in the image.

We further perform cross-verification by recoloring the mask contour and prompting the model to confirm that the expression is unique and unambiguous. This ensures high-quality referring segmentation annotations. To further enhance data diversity, we introduce reasoning segmentation. Unlike referring expressions that directly identify objects, reasoning expressions omit object names and instead describe them by function or commonsense attributes (e.g., “the furniture that can both sit and lie down” → sofa). This design encourages models to leverage semantic reasoning beyond surface lexical cues. This process differs in prompt design and omits the verification step.

In total, SA-OVRS contains 1.93M validated instance masks, 1.15M semantic samples, and 850K textual expressions (800K referring, 50K reasoning) spanning over 5,800 open-vocabulary labels. To further illustrate the open-vocabulary coverage of SA-OVRS, Figure 4 presents a word cloud of representative object categories appearing in our collected expressions.



Fig. 4: Visualizations of selected labels from the SA-OVRS dataset. The size of each label roughly reflects its abundance in the dataset.

5 Experiments

5.1 Experiment Settings

Training Data. For semantic segmentation, we select ADE20K [64] and COCO-Stuff [4] dataset. For referring segmentation, we choose the refCOCO [59] series and refCLEF [15] datasets. Additionally, we adopt the semantic and referring segmentation data from SA-OVRS.

Tasks and Evaluation Metrics. For closed-set and open-vocabulary semantic segmentation, we employed mean IoU (mIoU). For referring segmentation, we follow prior works and use the cumulative intersection over cumulative union (cIoU). To evaluate the contour accuracy of masks produced by visual generative models, we propose the average Hausdorff Distance (d_{AHD}) metric under multi-level IoU thresholds. Given the boundary point sets of a predicted mask and GT mask, d_{AHD} is defined as:

$$d_{AHD}(X, Y) = \frac{1}{2} \left(\frac{1}{X} \sum_{x \in X} \min_{y \in Y} dist(x, y) + \frac{1}{Y} \sum_{y \in Y} \min_{x \in X} dist(x, y) \right), \quad (3)$$

where $dist(x, y)$ denotes the Euclidean distance between x and y . This metric captures bidirectional distances between the two point sets, providing an aggregate measure of global boundary alignment, where a **lower score** indicates better performance. For each threshold in $[0.5, 0.6, 0.7, 0.8, 0.9]$, we collect predictions with IoU above it, compute d_{AHD} , and report the mean average Hausdorff distance ($mAHD$)

Pre-training. We adopt a two-stage training protocol consisting of pre-training followed by fine-tuning. During pre-training, the model is trained on the SA-OVRS dataset and the referring segmentation datasets. We use the AdamW [26] optimizer with a learning rate of 2×10^{-4} , $\beta_1 = 0.9$, $\beta_2 = 0.95$, weight decay 0.05, and run for 4 epochs.

Fine-tuning. The semantic and referring segmentation tasks are trained separately on their respective datasets. We use AdamW with a learning rate of 1×10^{-4} , $\beta_1 = 0.9$, $\beta_2 = 0.99$, zero weight decay, and a WarmupCosineDecay scheduler with 1% linear warmup. We fine-tune for 10 epochs on the semantic segmentation datasets and 20 epochs on the referring segmentation datasets. All training is carried out on 8 NVIDIA GPUs (A800 or H20) under PyTorch’s distributed data parallel framework.

Model Variants. For the LLaMA-based framework, mask tokens are generated via greedy search. When leveraging an MLLM-based framework, we apply LoRA [14] for efficient fine-tuning and remove the unconditional generation component while keeping other configurations unchanged. We develop two models within the LLaMA-based framework, namely **LlamaSeg-B (base)** and **LlamaSeg-L (large)**, and further construct **LlamaSeg-1B_(MLLM)** within the Pre-trained MLLM-based Segmentation Framework (with LoRA/Adapter).

Table 1: Comparisons with visual generative models and MLLMs on closed-set and open-vocabulary semantic segmentation datasets. Unified-IO and Unified-IO2 are evaluated using their open-source weights. “Params” denotes model size and “384 pix.” indicates a 384-pixel input resolution. **Bold** marks the best result.

Type	Model	Params	ADE20K	COCO-Stuff	PC-459	PC-59	PAS-20
MLLM	LaSagnA [49]	7B	42.0	43.9	9.8	39.6	61.8
Visual generative model	Unified-IO-Large	0.77B	39.9	50.7	-	62.8	62.7
	Unified-IO-XL	2.9B	45.2	55.2	-	64.0	74.9
	Unified-IO2-Large	1.1B	35.6	48.8	-	55.8	71.5
	Unified-IO2-XL	3.2B	39.4	49.9	-	56.9	71.5
	Unified-IO2-XXL	6.6B	51.9	55.1	-	-	-
Visual generative model (ours)	LlamaSeg-B	0.77B	50.5	54.5	28.5	61.7	74.5
	LlamaSeg-1B _(MLLM)	1B	45.1	50.1	35.6	58.2	71.1
	LlamaSeg-L	1.5B	52.0	55.7	35.5	62.5	75.1
	LlamaSeg-L _(384 pix.)	1.5B	55.9 (+3.9)	58.1 (+2.4)	36.7 (+1.2)	63.8 (+1.3)	77.1 (+2.0)

Table 2: Comparison results on referring segmentation datasets. We evaluated the performance of Unified-IO and Unified-IO2 on a subset of datasets. **Green** indicates the best result among discriminative segmentation models, while **orange** indicates the best result among visual generative models.

Type	Model	Params	refCOCO			refCOCO+			refCOCOg	
			val	testA	testB	val	testA	testB	val	test
Discriminative segmentation model	LAVT [56]	-	72.7	75.8	68.8	62.1	68.4	55.1	61.2	62.1
	ReLA [22]	-	73.8	76.5	70.2	66.0	71.0	57.7	65.0	66.0
Visual generative model	Unified-IO-Large	0.77B	35.7	39.2	-	34.7	39.4	-	42.2	42.3
	Unified-IO-XL	2.9B	42.4	49.5	-	39.4	42.7	-	50.3	50.3
	Unified-IO2-Large	1.1B	40.3	47.2	-	33.0	39.1	-	43.2	44.1
	Unified-IO2-XL	3.2B	50.7	55.7	-	39.4	47.4	-	52.6	54.5
	Unified-IO2-XXL	6.6B	54.8	-	-	44.8	-	-	-	-
Visual generative model (ours)	LlamaSeg-B	0.77B	50.9	56.1	38.9	38.9	44.5	32.9	51.0	50.5
	LlamaSeg-1B _(MLLM)	1B	52.3	57.8	47.1	44.6	51.6	36.5	49.0	49.7
	LlamaSeg-L	1.5B	56.5	60.5	52.6	41.6	46.2	36.2	51.0	50.4

5.2 Main Results

Semantic Segmentation Result Analysis. Table 1 reports closed-set and open-vocabulary semantic segmentation results. LlamaSeg consistently outperforms existing visual generative models across all benchmarks. LlamaSeg-L attains 52.0/55.7 mIoU on ADE20K/COCO-Stuff, surpassing Unified-IO-XL and Unified-IO2-XXL by large margins with fewer parameters. Using higher-resolution mask tokens, LlamaSeg-L (384 pix.) further improves to 55.9 (+3.9) and 58.1 (+2.4), with consistent gains on PC-459/PC-59/PAS-20 (e.g., +2.0 on PAS-20). These results show that our visual-token formulation scales well with model size and resolution, enabling fine-grained masks and better boundary preservation. Compared with the MLLM-based baseline (LaSagnA) and the Unified-IO

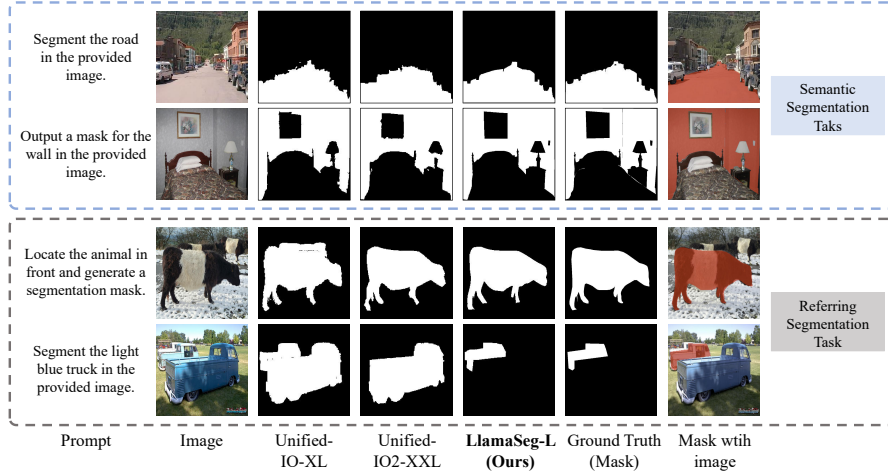


Fig. 5: Visual comparison between our model and other visual generative models.

Table 3: Comparison of the $mAHD$ metric after normalization to a 256-pixel resolution. Lower values correspond to a closer match between the predicted mask contours and the ground truth. Bold entries highlight the best results.

Model	ADE20K					refCOCO				
	IoU-0.5	IoU-0.6	IoU-0.7	IoU-0.8	IoU-0.9	IoU-0.5	IoU-0.6	IoU-0.7	IoU-0.8	IoU-0.9
Unified-IO-Base	76.41	74.84	72.36	68.72	62.39	85.14	83.22	81.26	80.02	79.06
Unified-IO-Large	76.35	75.31	72.21	67.73	62.62	83.08	81.02	79.81	78.38	75.61
Unified-IO-XL	75.88	75.16	72.35	69.35	65.13	72.58	71.69	71.16	69.99	85.65
Unified-IO2-Large	78.45	76.96	75.02	72.62	67.37	83.32	82.57	81.68	80.59	79.14
Unified-IO2-XL	80.15	78.51	76.48	74.21	67.82	82.05	81.42	80.91	79.95	78.18
Unified-IO2-XXL	79.72	78.02	75.76	71.81	65.96	80.98	80.10	79.45	78.24	75.33
LlamaSeg-B	26.61	25.91	25.61	25.98	29.28	14.40	13.63	12.69	11.36	9.94
LlamaSeg-1B _(MLLM)	59.24	58.91	59.20	60.77	69.38	45.47	42.41	39.29	37.80	37.28
LlamaSeg-L	25.54	24.73	24.28	24.40	26.51	10.45	9.68	8.73	7.42	5.27

series, LlamaSeg achieves higher accuracy while remaining fully autoregressive and end-to-end.

Referring Segmentation Result Analysis. As shown in Table 2, our model outperforms existing visual generative models across multiple referring segmentation datasets under the same parameter constraints. On the RefCOCO dataset, our LlamaSeg-Large achieves a score 1.7 points higher than the best result of Unified-IO2, demonstrating superior language understanding capabilities. In contrast, a noticeable gap remains compared to specialized discriminative segmentation models.

Boundary Accuracy Analysis. Although our model achieves slightly lower overall scores on referring segmentation tasks than other strong baselines, it de-

livers notably higher boundary precision. We evaluate $mAHD$ under multiple IoU thresholds on ADE20K and RefCOCO (Table 3). While semantic segmentation data, aimed at pixel-level discrimination, often exhibits complex contours, referring segmentation masks are comparatively simpler. Our approach consistently yields lower $mAHD$ than previous visual generative models, indicating more accurate alignment with ground-truth contours and the ability to capture fine object boundaries with high fidelity. Figure 5 further illustrates these improvements in mask quality and segmentation accuracy.

2D Spatial Relationship. To examine whether the model can learn 2D spatial relationships from 1D tokens, we select an example in resolution of 256 and export the attention map in the last self-attention layer of the Transformer model, as shown in Fig. 6. Under the causal attention mechanism, the attention map exhibits a distinct band parallel to the main diagonal, indicating that mask tokens preferentially attend to tokens in preceding rows within the same column in the original 2D layout. Moreover, we observe regular dark vertical stripes in the attention map. These stripes show that, although the row-end tokens and the next-row tokens are adjacent in the 1D sequence, the model consistently suppresses attention from the leading tokens of each row to the trailing tokens of the previous row. This behavior suggests that the model can distinguish true spatial neighbors in the 2D mask from artificial neighbors introduced by row-wise flattening. Therefore, these observations demonstrate that our method does not merely exploit local proximity in the 1D sequence. Instead, it internally reconstructs and leverages the authentic 2D spatial topology of mask tokens.

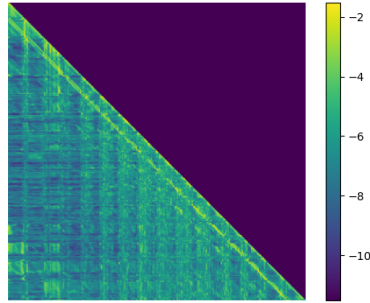


Fig. 6: Attention heatmap of mask tokens in the last self-attention layer of the autoregressive model. Scores have been log-transformed to enhance visual contrast.

5.3 Ablation Study

Mask Tokenizer Analysis. We further fine-tune the pre-trained VQGAN on segmentation masks with a batch size of 128. The overall IoU and $mAHD$ results (without IoU threshold constraints) for mask reconstruction are reported in Table 4. While incorporating additional training data slightly improves reconstruction quality, the overall performance remains below that of the original pre-trained model, suggesting that extended fine-tuning leads to a degradation of the learned feature representations. A pre-trained VQGAN, trained on numerous color images, has

Table 4: Ablation study on mask tokenizer. **Bold** indicates the best result.

Train steps	Total IoU↑	$mAHD$ ↓
0 (frozen)	96.8	6.1
5000	93.0	12.5
10000	93.5	10.6

strong visual feature-capturing and pixel reconstruction abilities. Since mask data only has black and white pixels, its features are simpler than those of color images, making it easier for the pre-trained VQGAN to capture edges. Although training the tokenizer on masked data seems to yield better results, it consumes significant resources. However, the model trained on large amounts of masked data still underperforms the frozen-parameter model. Visual examples are provided in Fig. 7

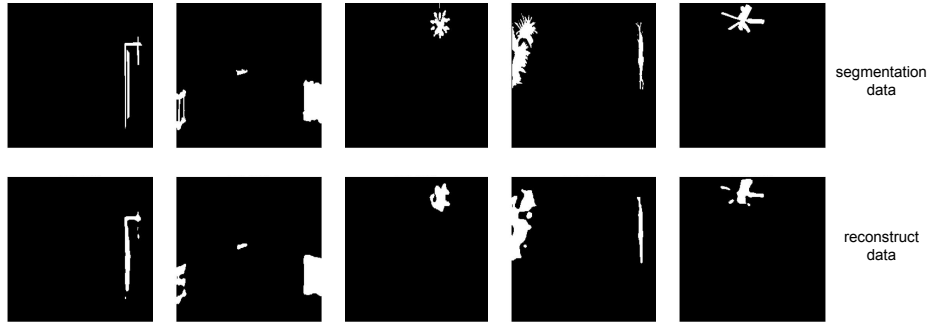


Fig. 7: Comparison between the segmentation masks and reconstructed masks. Although the IoU between the reconstructed and original mask is high, many detailed features have been lost.

Decoding Strategy Analysis. We conduct experiments on the RefCOCO validation set using the LlamaSeg-B model to systematically compare several decoding strategies, as reported in Table 5. Given that image segmentation tasks are paired with deterministic ground-truth masks, the model must produce a single, optimal prediction rather than diverse outputs. These results demonstrate that greedy decoding is the most effective and reliable strategy for autoregressive image segmentation in this setting.

Table 5: Ablation study on decoding strategy. The IoU threshold of mAHD is set to 0.5.

Decoding strategy	RefCOCO	
	cIoU $\uparrow$	mAHD $\downarrow$
Greedy search	50.9	14.4
Beam search (B=3)	47.3	15.1
Top-K (K=3)	49.4	15.3
Top-P (P=0.9)	50.2	15.3
Random sample	49.2	15.7

Table 6: Ablation study on training data. “Sem.” denotes semantic segmentation data and “Ref.” denotes referring segmentation data. The IoU threshold of mAHD is set to 0.5.

Pre-training data	Fine-tuning data	ADE20K		refCOCO	
		mIoU $\uparrow$	mAHD $\downarrow$	cIoU $\uparrow$	mAHD $\downarrow$
SA-OVRS+Ref.	Sem.	50.5	26.6	-	-
-	Sem.	48.9	28.7	-	-
SA-OVRS+Ref.	Ref.	-	-	50.9	14.4
-	Ref.	-	-	46.0	16.2

SA-OVRS Contribution Analysis. To evaluate the contribution of the SA-OVRS dataset to segmentation performance, we train LlamaSeg-B using different combinations of pre-training and fine-tuning data and summarize the results in Table 6. We compare models pre-trained with SA-OVRS+Ref. against counterparts trained without pre-training, while keeping the downstream fine-tuning setup identical for each task (ADE20K for semantic segmentation and refCOCO for referring segmentation). The ablation shows a clear performance drop when pre-training data are omitted, highlighting the importance of large-scale pre-training for capturing rich visual representations. Nevertheless, even without pre-training data, our model still surpasses existing visual generative models of comparable parameter size.

Efficiency Analysis. As shown in Table 7, LlamaSeg is slower than LAVT because autoregressive mask generation is sequential, whereas LAVT uses a parallel dense decoder. However, LlamaSeg is much faster than Unified-IO2-Large under the same evaluation protocol, demonstrating improved efficiency within the generative segmentation paradigm.

Table 7: Inference efficiency comparison on latency, FPS, and peak memory.

Method	Latency (ms)	FPS	Peak Memory (GB)
LAVT	35.89	27.84	1.78
Unified-IO2-Large	57906.23	0.017	3.93
LlamaSeg-B	3949.92	0.250	3.69

6 Conclusion

In this work, we propose **LlamaSeg**, a novel visual autoregressive image segmentation method that integrates segmentation into a standard autoregressive framework. We further introduce a two-stage data annotation pipeline to construct **SA-OVRS**, a large-scale open-vocabulary segmentation dataset containing 2M high-quality samples, which provides valuable resources for training and future research. Moreover, we introduce a new boundary-oriented evaluation metric $mAHD$ to more faithfully measure contour accuracy beyond region-level overlap metrics such as IoU and Dice. Extensive experiments show that LlamaSeg not only surpasses existing visual generative models of comparable size across multiple benchmarks but also provides a scalable and unified paradigm that can inspire future multimodal and autoregressive segmentation research.

7 Acknowledge

This research was financially supported by the Research Task Assignment Project from Guangdong Laboratory of Artificial Intelligence and Digital Economy (SZ), under Grant No.GML-26420006

References

- [1] Alayrac, J.B., Donahue, J., Luc, P., Miech, A., Barr, I., Hasson, Y., Lenc, K., Mensch, A., Millican, K., Reynolds, M., et al.: Flamingo: a visual language model for few-shot learning. *Advances in neural information processing systems* **35**, 23716–23736 (2022)
- [2] Bai, J., Bai, S., Yang, S., Wang, S., Tan, S., Wang, P., Lin, J., Zhou, C., Zhou, J.: Qwen-vl: A versatile vision-language model for understanding, localization, text reading, and beyond. *arXiv preprint arXiv:2308.12966* (2023)
- [3] Brown, T., Mann, B., Ryder, N., Subbiah, M., Kaplan, J.D., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., Askell, A., et al.: Language models are few-shot learners. *Advances in neural information processing systems* **33**, 1877–1901 (2020)
- [4] Caesar, H., Uijlings, J., Ferrari, V.: Coco-stuff: Thing and stuff classes in context. In: *Proceedings of the IEEE conference on computer vision and pattern recognition*. pp. 1209–1218 (2018)
- [5] Chen, J., Lu, J., Zhu, X., Zhang, L.: Generative semantic segmentation. In: *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*. pp. 7111–7120 (2023)
- [6] Chen, X., Wu, Z., Liu, X., Pan, Z., Liu, W., Xie, Z., Yu, X., Ruan, C.: Janus-pro: Unified multimodal understanding and generation with data and model scaling. *arXiv preprint arXiv:2501.17811* (2025)
- [7] Cheng, B., Schwing, A., Kirillov, A.: Per-pixel classification is not all you need for semantic segmentation. *Advances in neural information processing systems* **34**, 17864–17875 (2021)
- [8] Chowdhery, A., Narang, S., Devlin, J., Bosma, M., Mishra, G., Roberts, A., Barham, P., Chung, H.W., Sutton, C., Gehrmann, S., et al.: Palm: Scaling language modeling with pathways. *Journal of Machine Learning Research* **24**(240), 1–113 (2023)
- [9] Dai, W., Li, J., Li, D., Tiong, A., Zhao, J., Wang, W., Li, B., Fung, P., Hoi, S.: InstructBLIP: Towards general-purpose vision-language models with instruction tuning. In: *Thirty-seventh Conference on Neural Information Processing Systems* (2023), <https://openreview.net/forum?id=vvoWPYqZJA>
- [10] Devlin, J., Chang, M.W., Lee, K., Toutanova, K.: Bert: Pre-training of deep bidirectional transformers for language understanding. In: *Proceedings of the 2019 conference of the North American chapter of the association for computational linguistics: human language technologies, volume 1 (long and short papers)*. pp. 4171–4186 (2019)
- [11] Ding, H., Liu, C., Wang, S., Jiang, X.: Vision-language transformer and query generation for referring segmentation. In: *Proceedings of the IEEE/CVF international conference on computer vision*. pp. 16321–16330 (2021)
- [12] Esser, P., Rombach, R., Ommer, B.: Taming transformers for high-resolution image synthesis. In: *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*. pp. 12873–12883 (2021)

- [13] Hendrycks, D., Gimpel, K.: Gaussian error linear units (gelus). arXiv preprint arXiv:1606.08415 (2016)
- [14] Hu, E.J., Shen, Y., Wallis, P., Allen-Zhu, Z., Li, Y., Wang, S., Wang, L., Chen, W., et al.: Lora: Low-rank adaptation of large language models. ICLR **1**(2), 3 (2022)
- [15] Kazemzadeh, S., Ordonez, V., Matten, M., Berg, T.: Referitgame: Referring to objects in photographs of natural scenes. In: Proceedings of the 2014 conference on empirical methods in natural language processing (EMNLP). pp. 787–798 (2014)
- [16] Kingma, D.P., Welling, M., et al.: Auto-encoding variational bayes (2013)
- [17] Kirillov, A., Mintun, E., Ravi, N., Mao, H., Rolland, C., Gustafson, L., Xiao, T., Whitehead, S., Berg, A.C., Lo, W.Y., et al.: Segment anything. In: Proceedings of the IEEE/CVF international conference on computer vision. pp. 4015–4026 (2023)
- [18] Lai, X., Tian, Z., Chen, Y., Li, Y., Yuan, Y., Liu, S., Jia, J.: Lisa: Reasoning segmentation via large language model. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 9579–9589 (2024)
- [19] Lan, M., Chen, C., Zhou, Y., Xu, J., Ke, Y., Wang, X., Feng, L., Zhang, W.: Text4seg: Reimagining image segmentation as text generation. arXiv preprint arXiv:2410.09855 (2024)
- [20] Li, B., Zhang, Y., Guo, D., Zhang, R., Li, F., Zhang, H., Zhang, K., Zhang, P., Li, Y., Liu, Z., et al.: Llava-onevision: Easy visual task transfer. arXiv preprint arXiv:2408.03326 (2024)
- [21] Li, J., Li, D., Savarese, S., Hoi, S.: Blip-2: Bootstrapping language-image pre-training with frozen image encoders and large language models. In: International conference on machine learning. pp. 19730–19742. PMLR (2023)
- [22] Liu, C., Ding, H., Jiang, X.: Gres: Generalized referring expression segmentation. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 23592–23601 (2023)
- [23] Liu, H., Li, C., Li, Y., Li, B., Zhang, Y., Shen, S., Lee, Y.J.: Llava-next: Improved reasoning, ocr, and world knowledge (2024)
- [24] Liu, H., Li, C., Wu, Q., Lee, Y.J.: Visual instruction tuning. Advances in neural information processing systems **36**, 34892–34916 (2023)
- [25] Liu, S., Zeng, Z., Ren, T., Li, F., Zhang, H., Yang, J., Jiang, Q., Li, C., Yang, J., Su, H., et al.: Grounding dino: Marrying dino with grounded pre-training for open-set object detection. In: European Conference on Computer Vision. pp. 38–55. Springer (2024)
- [26] Loshchilov, I., Hutter, F.: Decoupled weight decay regularization. arXiv preprint arXiv:1711.05101 (2017)
- [27] Lu, J., Clark, C., Lee, S., Zhang, Z., Khosla, S., Marten, R., Hoiem, D., Kembhavi, A.: Unified-io 2: Scaling autoregressive multimodal models with vision language audio and action. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 26439–26455 (2024)
- [28] Lu, J., Clark, C., Zellers, R., Mottaghi, R., Kembhavi, A.: Unified-io: A unified model for vision, language, and multi-modal tasks. arXiv preprint arXiv:2206.08916 (2022)

- [29] Pramanick, S., Han, G., Hou, R., Nag, S., Lim, S.N., Ballas, N., Wang, Q., Chellappa, R., Almahairi, A.: Jack of all tasks master of many: Designing general-purpose coarse-to-fine vision-language model. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 14076–14088 (2024)
- [30] Radford, A., Narasimhan, K., Salimans, T., Sutskever, I., et al.: Improving language understanding by generative pre-training (2018)
- [31] Ramesh, A., Pavlov, M., Goh, G., Gray, S., Voss, C., Radford, A., Chen, M., Sutskever, I.: Zero-shot text-to-image generation. In: International conference on machine learning. pp. 8821–8831. Pmlr (2021)
- [32] Razavi, A., Van den Oord, A., Vinyals, O.: Generating diverse high-fidelity images with vq-vae-2. *Advances in neural information processing systems* **32** (2019)
- [33] Ren, Z., Huang, Z., Wei, Y., Zhao, Y., Fu, D., Feng, J., Jin, X.: Pixellm: Pixel reasoning with large multimodal model. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 26374–26383 (2024)
- [34] Singh, A., Rozeboom, J., Ray, N.: Tokenizing semantic segmentation with run length encoding. *arXiv preprint arXiv:2602.21627* (2026)
- [35] Sun, P., Jiang, Y., Chen, S., Zhang, S., Peng, B., Luo, P., Yuan, Z.: Autoregressive model beats diffusion: Llama for scalable image generation. *arXiv preprint arXiv:2406.06525* (2024)
- [36] Team, C.: Chameleon: Mixed-modal early-fusion foundation models, 2024. URL <https://arxiv.org/abs/2405.09818> **9**
- [37] Tian, K., Jiang, Y., Yuan, Z., Peng, B., Wang, L.: Visual autoregressive modeling: Scalable image generation via next-scale prediction. *Advances in neural information processing systems* **37**, 84839–84865 (2024)
- [38] Touvron, H., Lavril, T., Izacard, G., Martinet, X., Lachaux, M.A., Lacroix, T., Rozière, B., Goyal, N., Hambro, E., Azhar, F., et al.: Llama: Open and efficient foundation language models. *arXiv preprint arXiv:2302.13971* (2023)
- [39] Touvron, H., Martin, L., Stone, K., Albert, P., Almahairi, A., Babaei, Y., Bashlykov, N., Batra, S., Bhargava, P., Bhosale, S., et al.: Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288* (2023)
- [40] Tschannen, M., Gritsenko, A., Wang, X., Naeem, M.F., Alabdulmohsin, I., Parthasarathy, N., Evans, T., Beyer, L., Xia, Y., Mustafa, B., et al.: Siglip 2: Multilingual vision-language encoders with improved semantic understanding, localization, and dense features. *arXiv preprint arXiv:2502.14786* (2025)
- [41] Van Den Oord, A., Vinyals, O., et al.: Neural discrete representation learning. *Advances in neural information processing systems* **30** (2017)
- [42] Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A.N., Kaiser, Ł., Polosukhin, I.: Attention is all you need. *Advances in neural information processing systems* **30** (2017)

- [43] Wang, H., Tang, H., Jiang, L., Shi, S., Naeem, M.F., Li, H., Schiele, B., Wang, L.: Git: Towards generalist vision transformer through universal language interface. In: European Conference on Computer Vision. pp. 55–73. Springer (2024)
- [44] Wang, P., Bai, S., Tan, S., Wang, S., Fan, Z., Bai, J., Chen, K., Liu, X., Wang, J., Ge, W., et al.: Qwen2-vl: Enhancing vision-language model’s perception of the world at any resolution. arXiv preprint arXiv:2409.12191 (2024)
- [45] Wang, T., Cheng, C., Wang, L., Chen, S., Zhao, W.: Himtok: Learning hierarchical mask tokens for image segmentation with large multimodal model. In: Proceedings of the IEEE/CVF International Conference on Computer Vision. pp. 23267–23278 (2025)
- [46] Wang, W., Chen, Z., Chen, X., Wu, J., Zhu, X., Zeng, G., Luo, P., Lu, T., Zhou, J., Qiao, Y., et al.: Visionllm: Large language model is also an open-ended decoder for vision-centric tasks. Advances in Neural Information Processing Systems **36**, 61501–61513 (2023)
- [47] Wang, X., Ru, L., Huang, Z., Ji, K., Zheng, D., Chen, J., Zhou, J.: Argenseg: Image segmentation with autoregressive image generation model. Advances in Neural Information Processing Systems **38**, 28312–28336 (2026)
- [48] Wang, Z., Lu, Y., Li, Q., Tao, X., Guo, Y., Gong, M., Liu, T.: Cris: Clip-driven referring image segmentation. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 11686–11695 (2022)
- [49] Wei, C., Tan, H., Zhong, Y., Yang, Y., Lasagna, L.M.: Language-based segmentation assistant for complex queries. arXiv preprint arXiv:2404.08506 **2**(3), 6 (2024)
- [50] Weng, T., Jiang, W., Wang, J., Li, M., Ma, L., Ming, Z.: Oddgridbench: Exposing the lack of fine-grained visual discrepancy sensitivity in multimodal large language models. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 1575–1584 (June 2026)
- [51] Weng, T., Wang, J., Jiang, W., Ming, Z.: Visnumbench: Evaluating number sense of multimodal large language models. arXiv preprint arXiv:2503.14939 (2025)
- [52] Wu, J., Zhong, M., Xing, S., Lai, Z., Liu, Z., Chen, Z., Wang, W., Zhu, X., Lu, L., Lu, T., et al.: Visionllm v2: An end-to-end generalist multimodal large language model for hundreds of vision-language tasks. Advances in Neural Information Processing Systems **37**, 69925–69975 (2024)
- [53] Xia, Z., Han, D., Han, Y., Pan, X., Song, S., Huang, G.: Gsva: Generalized segmentation via multimodal large language models. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 3858–3869 (2024)
- [54] Xie, E., Wang, W., Yu, Z., Anandkumar, A., Alvarez, J.M., Luo, P.: Segformer: Simple and efficient design for semantic segmentation with transformers. Advances in neural information processing systems **34**, 12077–12090 (2021)

- [55] Xu, J., Xu, L., Yang, Y., Li, X., Wang, F., Xie, Y., Huang, Y.J., Li, Y.: u-llava: Unifying multi-modal tasks via large language model. In: ECAI 2024, pp. 618–625. IOS Press (2024)
- [56] Yang, Z., Wang, J., Tang, Y., Chen, K., Zhao, H., Torr, P.H.: Lavt: Language-aware vision transformer for referring image segmentation. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 18155–18165 (2022)
- [57] Yu, J., Li, X., Koh, J.Y., Zhang, H., Pang, R., Qin, J., Ku, A., Xu, Y., Baldrige, J., Wu, Y.: Vector-quantized image modeling with improved vq-gan. arXiv preprint arXiv:2110.04627 (2021)
- [58] Yu, J., Xu, Y., Koh, J.Y., Luong, T., Baid, G., Wang, Z., Vasudevan, V., Ku, A., Yang, Y., Ayan, B.K., et al.: Scaling autoregressive models for content-rich text-to-image generation. arXiv preprint arXiv:2206.10789 **2**(3), 5 (2022)
- [59] Yu, L., Poirson, P., Yang, S., Berg, A.C., Berg, T.L.: Modeling context in referring expressions. In: Computer Vision–ECCV 2016: 14th European Conference, Amsterdam, The Netherlands, October 11–14, 2016, Proceedings, Part II 14. pp. 69–85. Springer (2016)
- [60] Zhai, X., Mustafa, B., Kolesnikov, A., Beyer, L.: Sigmoid loss for language image pre-training. In: Proceedings of the IEEE/CVF international conference on computer vision. pp. 11975–11986 (2023)
- [61] Zhang, T., Li, X., Fei, H., Yuan, H., Wu, S., Ji, S., Loy, C.C., Yan, S.: Omg-llava: Bridging image-level, object-level, pixel-level reasoning and understanding. *Advances in Neural Information Processing Systems* **37**, 71737–71767 (2024)
- [62] Zheng, C., Vuong, T.L., Cai, J., Phung, D.: Movq: Modulating quantized vectors for high-fidelity image generation. *Advances in Neural Information Processing Systems* **35**, 23412–23425 (2022)
- [63] Zheng, R., Qi, L., Chen, X., Wang, Y., Wang, K., Zhao, H.: Seg-var: Image segmentation with visual autoregressive modeling. *Advances in Neural Information Processing Systems* **38**, 165921–165942 (2026)
- [64] Zhou, B., Zhao, H., Puig, X., Xiao, T., Fidler, S., Barriuso, A., Torrallba, A.: Semantic understanding of scenes through the ade20k dataset. *International Journal of Computer Vision* **127**, 302–321 (2019)