

# ViQ: Text-Aligned Visual Quantized Representations at Any Resolution

Xumin Yu<sup>1,\*\*</sup> , Zuyan Liu<sup>1,2,\*</sup> , Zhenyu Yang<sup>1,4,\*</sup>, Yuhao Dong<sup>3</sup>, Shengsheng Qian<sup>4</sup>, Jiwen Lu<sup>2†</sup> , Han Hu<sup>1</sup>, and Yongming Rao<sup>1†</sup> 

<sup>1</sup> Tencent HY Vision Team,

<sup>2</sup> Department of Automation, Tsinghua University,

<sup>3</sup> S-Lab, Nanyang Technological University,

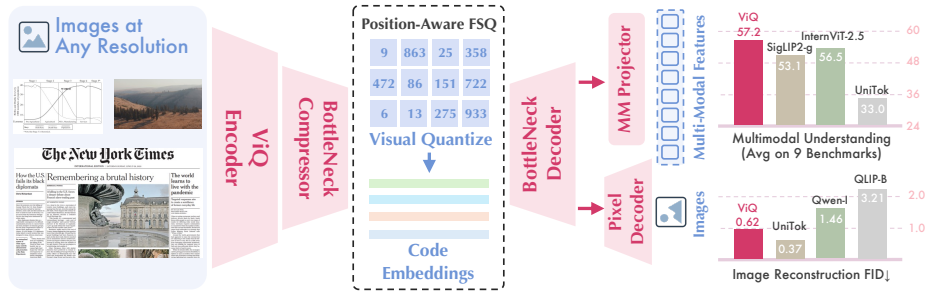
<sup>4</sup> Institute of Automation, Chinese Academy of Science

**Abstract.** A unified representation for text and vision is a natural pursuit, as it enables simpler multimodal modeling and more efficient training. However, representing images as discrete signals in the same way as text inevitably introduces severe information loss. Existing work struggles to balance low-level details and high-level semantics in discrete representations: reconstruction-oriented representations often lack semantic information, whereas semantically stronger features typically suffer from severe loss of detail. We present **ViQ**, a **V**isual **Q**uantized Representations framework, which is designed to balance semantics and details in discrete representations while supporting inputs at native resolutions, thereby enabling it to serve as a unified and general discrete representation for arbitrary visual inputs. Our approach structures quantization learning into two stages: text-aligned pre-training and feature discretization. With text-aligned pre-training, we enhance the visual encoder semantic-rich supervision from the pretrained language model and enable it to process native-resolution visual inputs. During discretization, we propose a proximal representation learning strategy to progressively compact the feature space, along with a position-aware head-wise quantization mechanism that enables flexible processing of arbitrary resolutions. Extensive experiments on multimodal tasks demonstrate that ViQ achieves competitive performance compared to state-of-the-art multimodal vision encoders with continuous and high-dimensional visual features, while maintaining high precision in low-level reconstruction. We also show that multimodal training with visual quantized representations largely improves efficiency, yielding up to 20%-70% acceleration with different base LLMs and training recipes.

**Keywords:** Visual Tokenization · Multimodal Representation Learning · Native Resolution

---

\* Equal Contribution. †Corresponding authors.



**Fig. 1: ViQ delivers high-quality multimodal quantized representations with both high-level semantics and low-level details.** The quantized visual codes in ViQ support high-level multimodal understanding and low-level image reconstruction, with state-of-the-art performance compared with continuous visual encoders.

## 1 Introduction

The rapid development of multimodal large language models (MLLMs) [4, 31, 52] has created a growing demand for high-quality, unified visual representations. A key challenge in this field lies in aligning visual signals into a form that can be seamlessly interpreted by language models. Early and dominant approaches have primarily relied on continuous visual encoders [10, 32, 49], pre-trained through contrastive vision-language learning, or specialized encoders fine-tuned for multimodal tasks [3, 22, 31]. Although this paradigm has achieved considerable success, it introduces a fundamental representational discrepancy: the continuous nature of visual features is intrinsically mismatched with the discrete token-based modeling of text. Furthermore, the computational process of extracting high-dimensional visual features places significant hardware strain during model training.

Inspired by the success of vector quantization [11, 29, 47] in visual representation learning, a growing body of research has begun exploring discrete visual representations in multimodal domains [25, 38, 51]. These approaches typically quantize visual inputs into a finite set of discrete tokens, analogous to words in a textual vocabulary. However, a fundamental challenge remains in balancing low-level visual details with high-level semantic information. Reconstruction-focused autoencoders often fail to preserve semantic structure in their latent features, while semantically rich visual encoders tend to incur significant information loss during quantization. This creates a critical gap in multimodal representations, which require both fine-grained visual details and rich semantic content to handle complex tasks effectively. As a result, the application of quantized visual encoders has been largely limited, and they have yet to match the high fidelity and robustness of continuous encoders in demanding real-world scenarios.

To address this issue, we propose ViQ (Visual Quantized Representations), a novel approach that elevates the performance of discrete visual representations in multimodal tasks to a level comparable with widely-used continuous features, while supporting native-resolution visual inputs. Our method structures quanti-

zation learning into two phases: text-aligned pre-training and visual quantization learning. In the initial stage, the ViQ model learns semantically rich alignments through supervision from vision-language pairs. To mitigate information loss during quantization, we introduce a proximal representation learning strategy that constrains the latent visual space, followed by a quantization mechanism enhanced with position encoding and expanded visual features. This design supports quantization at arbitrary resolutions and improves the representational capacity of visual codes. Our work not only strongly enhances the efficiency of multimodal encoders in training process but also unifies visual and linguistic representations into a cohesive discrete framework.

We evaluate the ViQ model against state-of-the-art continuous multimodal visual encoders and quantized multimodal semantic encoders across a range of experiments. Under consistent fine-tuning data and training protocols, ViQ not only outperforms existing quantized models by a large margin but also achieves competitive results compared to well-established continuous encoders such as InternViT [52], AIMv2 [10], and SigLIP2 [40]. On the aggregated score over nine benchmarks—spanning visual question answering, world knowledge, and document and chart recognition—ViQ attains an average of 57.2 with Qwen2.5-1.5B as the backbone LLM and 63.9 with Qwen2.5-7B, surpassing the previous state-of-the-art scores of 57.0 and 63.8, respectively, under 6B number of visual encoder parameters. The quantized representation also brings substantial efficiency gains in real-world multimodal training. Compared to conventional training strategies, using the ViQ visual encoder yields speedups of 20% to 70% in training recipes varying in sequence lengths. Furthermore, ViQ preserves rich low-level visual details: when fine-tuned with an image decoder, it achieves a PSNR of 22.73 and an rFID score of 0.62, ranking first among mainstream discrete visual autoencoders. This strong performance in both understanding and reconstruction tasks underscores the effectiveness and unity of our discrete representation.

## 2 Related Works

**Visual Representations for MLLMs.** The visual encoder serves as a critical bridge between vision and language modalities in multimodal learning. Conventional multimodal language models typically employ CLIP-style models—such as CLIP [32], SigLIP [49], SigLIP2 [40], AIM [10], and DFN [9]—as visual encoders. However, such models often rely on fixed input resolutions, which constrains their flexibility in multimodal tasks, and their contrastively learned representations may not align optimally with downstream multimodal objectives. To address these limitations, several specialized MLLM visual encoders have been developed and trained end-to-end on multimodal tasks. These include open-source models like InternViT [4, 52], OryxViT [22], and SAILViT [46], as well as proprietary encoders integrated into large MLLM systems such as Qwen-VL [39], Kimi-VL [37], and Seed-VL [13]. More recently, researchers have begun exploring unified discrete representations for vision and language. Quantized multimodal encoders such

as CLIP [51] and UniTok [25] apply quantization to image inputs. Nevertheless, these quantized visual encoders still exhibit a large performance gap compared to continuous models, particularly in tasks requiring textual understanding or fine-grained visual details. This gap can be attributed to the high compression ratio of discrete codes and the high sensitivity of multimodal models to detailed visual information.

**Quantized Visual Encoders.** Image tokenization is pivotal in bridging raw pixels with compact latent representations for visual generation. Among existing approaches, vector-quantized (VQ) [11] tokenizers have gained prominence due to their discrete latent space, which is well-suited for visual generation. The seminal VQ-VAE [41] introduced a learnable codebook to discretize continuous features via nearest-neighbor assignment. Subsequent methods (such as FSQ [29], RPQ [6], BSQ [50], and LFQ [47]) explored structurally constrained codebooks with predefined geometries, often decoupled from end-to-end training. Enhancements like VQ-GAN [8] further improved visual quality by incorporating adversarial and perceptual losses. While these tokenizers excel at learning compact, generative-friendly representations, they often struggle to preserve fine-grained visual details crucial for dense prediction and high-fidelity understanding tasks. To address this, we propose ViQ, a novel framework designed to minimize information loss and align discrete visual tokens with semantic and textual contexts.

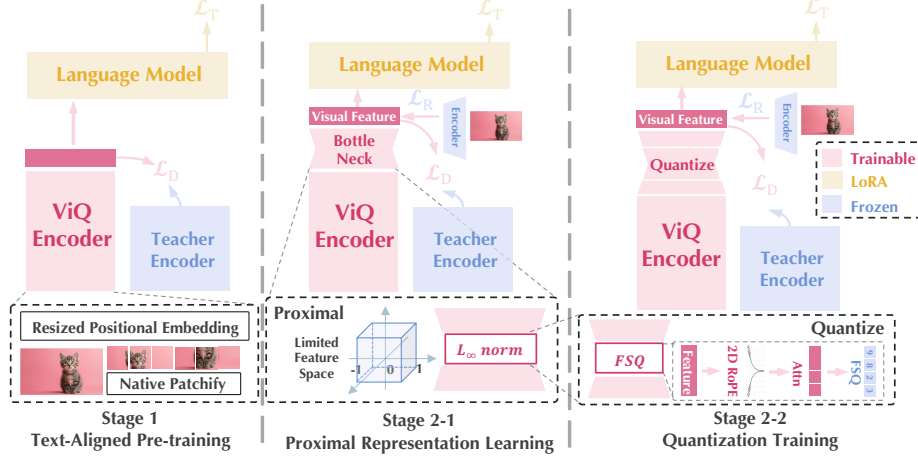
### 3 ViQ: Visual Quantized Representations

The visual quantization training for ViQ follows a two-stage process, as depicted in Fig. 2. In Stage 1, text-aligned pre-training is performed on continuous features to enhance multimodal alignment. Subsequently, Stage 2 applies quantization in a progressive manner.

#### 3.1 Text-Aligned Pre-Training at Any Resolution

The first stage of ViQ training aims to create a visual encoder that functions in a multi-modal manner. This is accomplished by leveraging text-aligned visual pre-training to align the vision and language embeddings within ViQ, thereby following an optimization strategy similar to conventional multi-modal training.

**Any Resolution Adaptation.** While Language-Image Pre-training models such as CLIP [32] and SigLIP [49] offer a strong initialization for building multi-modal visual encoders, most existing models are constrained by a fixed input size. In multi-modal learning, native image resolution often provides a more natural and efficient form of visual representation. To overcome the limitations of fixed-scale pre-training, we follow the approach of NaViT [7] and replace the original positional embedding layer with one that supports an adequately dimensioned size. This allows the model to resize positional parameters when processing images at arbitrary resolutions dynamically. We adopt the techniques introduced in OryxViT [22] to maintain computational efficiency with variable-length visual tokens.



**Fig. 2: Approach of ViQ Representation Learning.** Stage 1 enables multimodal alignment with language supervision, while Stage 2 compresses the high-dimensional visual features into discrete codes in a progressive learning manner.

**Text-Guided Multi-Modal Pre-training.** We optimize the continuous feature space of ViQ to be more compatible with multi-modal learning during our text-aligned pre-training stage. Specifically, given a data triplet  $[I, T, A]$  consisting of an image  $I$ , a text query  $T$ , and an answer  $A$ , we employ a temporary language model to compute the text supervision loss  $\mathcal{L}_{\text{text}}$ , which is defined as:

$$\mathcal{L}_{\text{text}} = \text{Cross Entropy}[\text{LLM}(\text{ViQ}(I), T), A] \quad (1)$$

**Self Distillation.** During the multi-modal pre-training stage, we apply self-distillation to regularize the ViQ model, preventing it from overfitting to the multi-modal data and disregarding the knowledge acquired from large-scale language-image pre-training, which is crucial for maintaining the generalizability and foundational capabilities. We use the original fixed-resolution model as the teacher to supervise the cosine similarity of the semantic token (i.e., the class token in mainstream architectures), ensuring that the representations produced at any resolution preserve the original high-level semantic information:

$$\mathcal{L}_{\text{distill}} = 1 - \cos(\mathbf{z}_s^{\text{student}}, \mathbf{z}_s^{\text{teacher}}) \quad (2)$$

where  $\mathbf{z}_s^{\text{student}}$  and  $\mathbf{z}_s^{\text{teacher}}$  denote the semantic tokens from the student (any-resolution) and teacher (fixed-resolution) models, respectively.

**Training Recipe for Native Resolution.** We employ a progressive training strategy that transitions from fixed-resolution to any-resolution inputs. Specifically, the training begins with low-resolution images at their native aspect ratios, and the total number of pixels is gradually increased throughout the training process until native resolution is attained. During training, multi-modal data are

packed to fit the maximum token length for higher pre-training efficiency. The image complexity, question-answering difficulty, and overall sequence length are progressively raised to steadily adapt the model to multi-modal tasks. Further implementation details are provided in the appendix. We combine the text loss  $\mathcal{L}_{\text{text}}$  and the distillation loss  $\mathcal{L}_{\text{distill}}$  for joint optimization.

### 3.2 Visual Quantized Representation Learning

The second stage of ViQ training involves quantizing the continuous features into discrete codes. A core challenge in this process is to fully optimize the discrete representation space so as to preserve fine-grained visual information. Through our approach, we demonstrate that the resulting ViQ representations are effective in supporting high-quality visual understanding in multimodal contexts.

**Proximal Representation.** The primary objective of quantization is to map continuous visual features from a high-dimensional space  $\mathbb{R}^C$  into a constrained discrete space defined by a codebook  $\mathcal{Z}$ . However, we observe that directly quantizing high-dimensional visual features results in significant precision loss. To mitigate this, our proposed ViQ method progressively reduces the complexity of the latent space using a proximal representation. Given a high-dimensional visual feature  $f \in \mathbb{R}^C$ , we first apply a bottleneck layer to compress it into an intermediate-dimensional feature  $f_1 \in \mathbb{R}^D$ . We then further constrain the feature space complexity via a regularization function, referred to as a proximal representation. In our implementation, we apply the  $L_\infty$ -norm to the compressed latent feature, projecting all features onto a hypercube surface such that  $\|f_1\|_\infty = 1$ . This step progressively reduces the feature space complexity prior to discrete quantization. The proximal representation helps regulate the distance between features and quantization anchors, thereby reducing information loss during quantization. The bottleneck compression operations can be formulated as:

$$f_1 = L_\infty(\text{BN}(f)), \hat{f} = \text{BN}'(f_1) \quad (3)$$

where BN denotes the bottleneck fully connected layer for dimension compression and  $\text{BN}'$  denotes the inverted bottleneck layer respectively.

**Multi-Head Finite Scalar Quantization.** After learning an initial continuous proximal representation within a constrained latent space, we replace the normalization function with a bottleneck downsampling layer that projects the latent features into a lower-dimensional space  $f_2 \in \mathbb{R}^d$ , where  $d \ll D \ll C$ . The final dimension  $d$  is kept sufficiently small to facilitate effective quantization. We employ Finite Scalar Quantization (FSQ) [29] to discretize the compact continuous features  $f_2$ , as this method offers stable training behavior without requiring additional optimization during quantization. The FSQ process can be formulated as:

$$z = \text{round}(\mathcal{Q}(f_2)) \quad (4)$$

where  $\mathcal{Q}$  indicates the Finite Scalar Quantization function.

To enhance the representational capacity of the quantized codes, we utilize a multi-head attention mechanism that expands each visual patch into  $2 \times 2$  visual

codes. Specifically, the latent feature  $f_2 \in \mathbb{R}^{B \times N \times d}$  is up-projected to  $\mathbb{R}^{B \times 4N \times d}$ , where  $N$  denotes the number of original image patches. Here the up-projection is a linear layer  $\mathbb{R}^d \rightarrow \mathbb{R}^{4d}$  that lifts each patch and reshapes it into 4 sub-tokens. The multi-head self-attention is applied only *within* the 4 sub-tokens of the same patch, so that different patches never interact; quantization is then performed element-wise on these tokens. Finally, the 4 quantized sub-tokens of each patch are concatenated and passed through a linear projection  $\mathbb{R}^{4d} \rightarrow \mathbb{R}^d$  to restore one token per patch and maintain the original downsampling rate. It is important to note that the expanded image patches are processed independently; as a result, the quantized representations in ViQ remain independent, making them more suitable for representation learning and downstream tasks.

**Rotary Position Embedding.** For quantized representations at arbitrary resolutions, we incorporate a 2D rotary position embedding (RoPE) [35] to explicitly encode spatial resolution information. The RoPE layer is inserted prior to the quantization step. Specifically, 2D RoPE extends the 1D RoPE formulation used in large language models by embedding both height and width information of the visual content. Given a token feature  $f_m \in \mathbb{R}^d$  at position  $(h, w)$  in the 2D feature map, the rotary encoding is applied as:

$$\tilde{f}_m = f_m \odot e^{i(h\theta_h + w\theta_w)} \quad (5)$$

where  $\theta_h$  and  $\theta_w$  are frequency parameters for the height and width dimensions, respectively, and  $\odot$  denotes element-wise multiplication with complex exponentials. This formulation ensures that relative spatial relationships are preserved across varying resolutions.

**Multi-Stage Training with Low-Level Supervision.** Building upon the progressive quantization framework, we learn the compressed feature representation in a multi-stage manner. First, starting from the proximal representation in the constrained space, we introduce a bottleneck layer with a regularization function between the continuous input features and the final output. We retain both the text loss  $\mathcal{L}_{\text{text}}$  and the self-distillation loss  $\mathcal{L}_{\text{distill}}$  during this phase, keeping all other settings consistent with Stage 1 training. To enhance low-level detail preservation and improve latent feature quality, we incorporate a pre-trained visual autoencoder that encodes the original input images into the low-level latent features. This encoding is supervised using a VAE-style latent loss along with a negative log-likelihood (NLL) term, expressed as:

$$\mathcal{L}_{\text{recon}} = \text{NLL}(\hat{f}, \text{Encoder}(x)) \quad (6)$$

where  $x$  denotes the input image and  $\hat{f}$  denotes the recovered features after compression or quantization. Here the NLL term is computed under a Gaussian likelihood with fixed (unit) variance over the target VAE latent  $\text{Encoder}(x)$ , so that the loss reduces, up to a constant, to a mean-squared error between  $\hat{f}$  and  $\text{Encoder}(x)$ , i.e.,  $\mathcal{L}_{\text{recon}} = \frac{1}{2} \|\hat{f} - \text{Encoder}(x)\|_2^2 + \text{const}$ . This makes the objective a simple and stable regression on the pre-trained VAE latent space rather than a pixel-level reconstruction.

Once the bottleneck compression is adequately learned, we replace the proximal regularization function with a quantization module in the subsequent quantization training stage. Since our vector quantization (VQ) implementation is optimization-free, no additional codebook supervision is required at this stage. The overall training objective combines the three losses as follows:

$$\mathcal{L}_{\text{total}} = \lambda_{\text{text}}\mathcal{L}_{\text{text}} + \lambda_{\text{distill}}\mathcal{L}_{\text{distill}} + \lambda_{\text{recon}}\mathcal{L}_{\text{recon}} \quad (7)$$

## 4 Experiments

In our experiments, we comprehensively evaluate the effectiveness of the ViQ model. We first describe the implementation details of ViQ’s key components. We then perform fair comparisons on multimodal tasks under consistent fine-tuning data and training protocols, benchmarking against leading multimodal and quantized encoders. To highlight practical benefits, we assess efficiency in real-world training and visual storage scenarios. We further examine the reconstruction quality to analyze low-level feature preservation, and provide ablation studies to validate design choices.

### 4.1 Implementation Details

We instantiate our ViQ model based on the pretrained SigLIP2-g [40] visual encoder, whose output feature dimension is  $C = 1536$ . To meet quantization requirements, we reduce the dimensionality through bottleneck layers to  $D = 128$  and subsequently to  $d = 6$ . For quantization, we employ the FSQ method [29] with levels  $\mathcal{L} = [8, 8, 8, 5, 5, 5]$ , yielding a codebook size of 64,000. Multimodal features from ViQ are projected via a simple MLP layer. During training, we use Qwen2.5-VL-0.5B as the language model for text supervision and the pretrained Qwen-Image [43] encoder for low-level visual supervision, keeping its parameters fixed. Stage 1 pretraining is conducted using 128 NVIDIA A100 GPUs, while Stage 2 quantization training uses 256 NVIDIA A100 GPUs. Further details regarding training stages and datasets are provided in the appendix.

### 4.2 Multi-Modal Understanding

In this subsection, we present experiments on multi-modal understanding to evaluate the effectiveness of our ViQ models in learning superior visual representations. By comparing ViQ with various baseline models, we show that our approach achieves compact visual representations without compromising their quality.

*Setups.* We integrate ViQ with large language models of varying scales to evaluate the generalization capability of our vision encoders. To compare with other visual encoders, as ViQ supports native-resolution perception, we adapt the LLaVA-NeXT [19]’s “any resolution” training pipeline for other visual encoders that



**Table 1: Comparison results on multi-modal understanding tasks.** We conduct experiments with various visual encoder counterparts on general multi-modal benchmarks. All the methods are trained with the same data collection.

| Base LLM             |      |        |          | MMStar | MMMU | SimpleVQA | InfoVQA | TextVQA | DocVQA | OCRBench | A12D | ChartQA | Avg. |
|----------------------|------|--------|----------|--------|------|-----------|---------|---------|--------|----------|------|---------|------|
| Visual Encoder       | Size | AnyRes | Discrete |        |      |           |         |         |        |          |      |         |      |
| Qwen2.5 - 1.5B       |      |        |          |        |      |           |         |         |        |          |      |         |      |
| OAI-CLIP-L [32]      | 0.3B | ✗      | ✗        | 44.9   | 40.7 | 21.3      | 24.2    | 58.9    | 52.9   | 460.0    | 68.1 | 55.0    | 45.8 |
| SigLIP2-g [40]       | 1.1B | ✗      | ✗        | 48.1   | 42.4 | 25.6      | 28.2    | 73.1    | 67.8   | 590.0    | 71.5 | 62.0    | 53.1 |
| DinoV2-g [30]        | 1.1B | ✗      | ✗        | 47.1   | 41.8 | 24.0      | 31.7    | 72.1    | 76.9   | 619.0    | 68.8 | 61.8    | 54.0 |
| OryxViT [22]         | 0.4B | ✓      | ✗        | 46.4   | 42.1 | 23.2      | 31.8    | 71.8    | 73.5   | 622.0    | 68.2 | 62.1    | 53.4 |
| AIMv2-H [10]         | 0.7B | ✗      | ✗        | 48.5   | 41.8 | 23.5      | 31.9    | 71.6    | 73.7   | 623.0    | 69.8 | 62.5    | 53.9 |
| InternViT-2.5 [3]    | 0.3B | ✗      | ✗        | 47.9   | 40.3 | 23.6      | 35.5    | 73.0    | 81.7   | 681.0    | 69.6 | 69.2    | 56.5 |
| InternViT-2.5-6B [3] | 6.0B | ✗      | ✗        | 48.5   | 42.1 | 23.7      | 35.2    | 75.5    | 80.1   | 690.0    | 70.7 | 67.8    | 57.0 |
| QLIP [51]            | 0.3B | ✗      | ✓        | 39.9   | 36.9 | 13.7      | 14.8    | 45.1    | 12.2   | 290.0    | 61.9 | 14.1    | 29.7 |
| UniTok [25]          | 0.3B | ✗      | ✓        | 41.0   | 36.1 | 15.5      | 15.9    | 39.7    | 11.6   | 323.0    | 61.2 | 43.8    | 33.0 |
| ViQ                  | 1.3B | ✓      | ✓        | 47.8   | 42.6 | 26.0      | 41.6    | 74.3    | 84.2   | 636.0    | 69.7 | 65.2    | 57.2 |
| Qwen2.5 - 7B         |      |        |          |        |      |           |         |         |        |          |      |         |      |
| OAI-CLIP-L [32]      | 0.3B | ✗      | ✗        | 53.9   | 47.1 | 25.4      | 33.9    | 66.4    | 61.4   | 544.0    | 76.6 | 65.1    | 53.8 |
| SigLIP2-g [40]       | 1.1B | ✗      | ✗        | 57.2   | 48.3 | 28.5      | 37.3    | 78.7    | 75.0   | 671.0    | 79.5 | 72.8    | 60.5 |
| OryxViT [22]         | 0.4B | ✓      | ✗        | 56.4   | 48.1 | 26.5      | 39.9    | 78.5    | 79.8   | 660.0    | 78.2 | 72.1    | 60.6 |
| AIMv2-H [10]         | 0.7B | ✗      | ✗        | 55.2   | 48.2 | 26.8      | 41.8    | 79.1    | 80.1   | 687.0    | 77.8 | 72.5    | 61.1 |
| InternViT-2.5-6B [3] | 6.0B | ✗      | ✗        | 55.3   | 48.1 | 28.4      | 44.9    | 79.9    | 85.7   | 757.0    | 78.7 | 77.4    | 63.8 |
| ViQ                  | 1.3B | ✓      | ✓        | 54.2   | 49.1 | 28.5      | 55.3    | 78.5    | 88.9   | 711.0    | 76.7 | 72.8    | 63.9 |

operate at fixed resolutions to ensure a fair comparison. For evaluation, we adopt representative visual understanding benchmarks, including general multi-modal benchmarks MMStar [2], MMMU [48], world knowledge-relevant benchmarks SimpleVQA [5], InfoVQA [27], text and doc recognition benchmarks including TextVQA [34], DocVQA [28], OCRBench [21], chart and scientific recognition benchmarks including A12D [14] and ChartQA [26]. The comprehensive benchmarks cover a broad range of visual skills, including basic perception, diagram interpretation, OCR, and visual reasoning. Each experiment is conducted on a fixed dataset of 2000K samples drawn from LLaVA-OneVision [18] to maintain consistency and fairness. We employ LMMs-Eval [17] as our evaluation toolkit to ensure reproducible results.

*Visual Encoder Baselines.* To enable a comprehensive comparison with existing visual encoders, we carefully select a diverse set of models across different categories. For general multi-modal visual encoders, we include widely used models such as OpenAI CLIP [32], SigLIP2 [40], and DINOv2 [30]. For visual encoders specialized optimized for mutli-modal data and tasks, we incorporate AIMv2 [10], OryxViT [22], and InternViT [52] (including 300M and 6B variants). In addition, we evaluate quantized visual encoders, including QLIP [51] and UniTok [25], to ensure a thorough and balanced comparison.

*Results.* As shown in Table 1, ViQ achieves competitive performance compared to other visual encoders. Despite being quantized, ViQ matches or surpasses representative general-purpose encoders on the aggregated score, primarily due to its native-resolution perception capability and quantization-aware training

strategy, which together preserve strong visual perception throughout quantization. Compared to multimodal encoders that are specifically trained for visual understanding, ViQ remains highly competitive overall and is particularly strong on text- and document-centric tasks. We note, however, that the advantage is not uniform: on certain detail-intensive benchmarks such as OCRBench, ViQ still trails some continuous encoders with fewer parameters, which we attribute to the inherent loss of high-frequency details when the continuous feature space is aggressively compressed into discrete codes. Moreover, ViQ consistently outperforms previous quantized encoders by a large margin, which typically suffer from severe degradation after quantization, highlighting the effectiveness of our proximal representation learning and quantization training strategy in preserving visual quality.

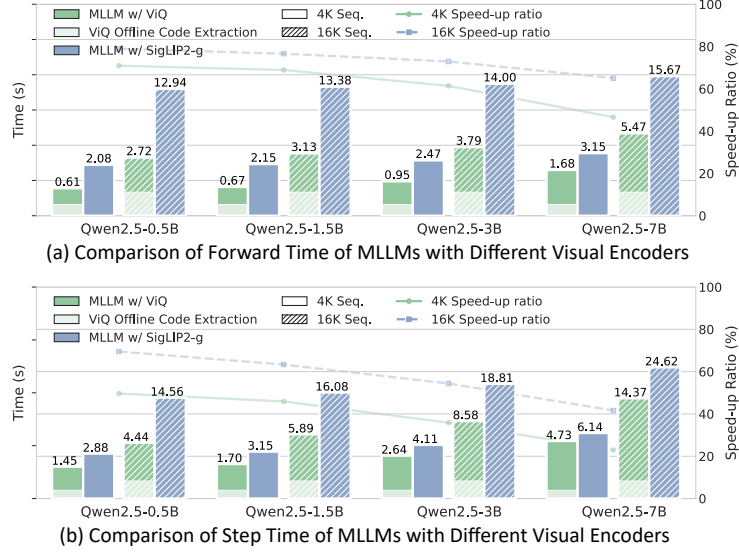
We further observe that ViQ’s gains are most pronounced on OCR, document, and infographic understanding tasks. We believe this is because general benchmarks depend more heavily on the backbone LLM’s knowledge and reasoning ability, whereas OCR-, document-, and chart-oriented tasks require precise low-level visual details and thus more directly reflect the capability of the visual encoder. From this perspective, the most meaningful contribution of ViQ is that it preserves such fine-grained understanding even after aggressively compressing continuous images into discrete tokens. The residual gap on the most detail-intensive tasks is a systemic property of discrete tokenization rather than a flaw specific to ViQ, and could be further narrowed by orthogonal directions such as multi-scale or residual quantization and incorporating specialized document data. Overall, ViQ serves as an efficient visual encoder that combines high compression with strong perceptual capability, making it well-suited for multimodal understanding tasks.

### 4.3 Efficiency

In this section, we demonstrate ViQ’s capabilities beyond multi-modal understanding through some experiments.

**Training Speed-Up for VLMs Setup.** We integrate ViQ and the popular SigLIP2-g encoder with a series of Qwen2.5 models for VLM SFT as conducted in LLaVA [20]. All experiments are conducted on a single node to eliminate noise introduced by inter-node communication. We fix the maximum image area to  $768^2$ , while preserving the original aspect ratio. We cover the 4k and 16k training cases in experiments, where image-text pair are packed into a fixed sequence length for stable training.

For SigLIP2-g, following common VLM training practice, we extract features from a list of images and feed them into the LLM together with processed text token embeddings for the forward pass. For ViQ, we first extract the discrete codes of each image offline. During the training, instead of loading raw images, we load the precomputed ViQ codes and project them into the LLM latent space.

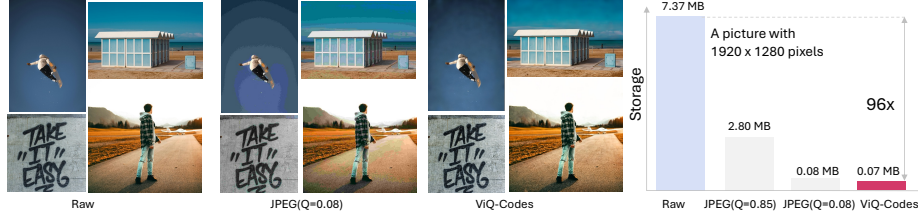


**Fig. 3: Comparisons on Training Efficiency Across Different Visual Encoders.** We conduct the experiments to compare the efficiency of ViQ and SigLIP2-g on the 4k and 16k training.

We discard the first 100 warm-up steps and report the average time required for a single forward pass and for a full training iteration.

**Results.** As shown in Figure 3, we compare the forward time and step time of SigLIP2-g and ViQ across four different Qwen2.5 model sizes, ranging from 0.5B to 7B. For a fair comparison, we also consider the ViQ offline code extraction time. We can see ViQ provides substantial speed-ups across all model sizes. The gains are especially pronounced for smaller LLMs: for the 0.5B model, ViQ accelerates forward time by 70% and 78% under the 4k and 16k training settings, respectively. And for larger model, like Qwen2.5 7B, ViQ can achieve a 46% and 65% speed-up in forward time under the 4k and 16k settings. Considering a whole iteration step, ViQ offers consistent improvements, exceeding 20% and 40% speed-ups in the 4k and 16k settings.

**Storing Any Image as Discrete Codes** As ViQ is a quantized visual encoder, it can convert an image at its native resolution into a series of discrete codes, which can then be reconstructed by a decoder. For an image of shape  $(H, W)$  with three channels, storing the raw image consumes  $\frac{H \times W \times 3 \times 8\text{bits}}{8} = 3HW$  bytes on disk. In contrast, ViQ can translate an image with  $(H, W)$  into  $\frac{H \times W}{64}$  codes, each ranging from 0 to 64,000, which can be represented using an unsigned 16-bit integer. As a result, storing the ViQ codes requires  $\frac{\frac{H \times W}{64} \times 16\text{bits}}{8} = \frac{HW}{32}$  bytes, which is only  $\frac{1}{96}$  of the raw image size. This corresponds to a very low target bitrate: matching the same compression ratio with JPEG would require an



**Fig. 4: Representing images with ViQ.** We show the image compression and visual reconstruction capability of ViQ. ViQ achieves a high-compression-ratio in image storage with high-quality reconstructed images while supporting native-resolution inputs.

**Table 2: Comparison of reconstruction quality on the  $256 \times 256$  ImageNet-1K validation set.** "Und." indicates whether the tokenizer is optimized for understanding tasks. "\*" means reproduced results with the official weights and the same evaluation code as our model.

| Method                      | Und.         | #Token         | PSNR $\uparrow$ | SSIM $\uparrow$ | rFID $\downarrow$ |
|-----------------------------|--------------|----------------|-----------------|-----------------|-------------------|
| <i>Continuous Tokenizer</i> |              |                |                 |                 |                   |
| SD-VAE 3                    | $\times$     | $32 \times 32$ | 31.29           | 0.87            | 0.20              |
| FLUX-VAE [16]               | $\times$     | $32 \times 32$ | 32.74           | 0.92            | 0.18              |
| Qwen-Image [43]             | $\times$     | $32 \times 32$ | 32.18           | 0.90            | 1.46              |
| Cosmos-CI [1]               | $\times$     | $16 \times 16$ | 25.07           | 0.70            | 0.96              |
| Wan2.2 [42]                 | $\times$     | $16 \times 16$ | 31.25           | 0.88            | 0.75              |
| <i>Discrete Tokenizer</i>   |              |                |                 |                 |                   |
| Cosmos-DI [1]               | $\times$     | $16 \times 16$ | 19.98           | 0.54            | 4.40              |
| Show-o [44]                 | $\times$     | $16 \times 16$ | 21.34           | 0.59            | 3.50              |
| LlamaGen [36]               | $\times$     | $16 \times 16$ | 20.65           | 0.54            | 2.47              |
| MUSE-VL [45]                | $\times$     | $16 \times 16$ | 20.14           | 0.65            | 2.26              |
| Open-MAGVIT2 [24]           | $\times$     | $16 \times 16$ | 22.70           | 0.64            | 1.67              |
| QLIP-B [51]                 | $\checkmark$ | $16 \times 16$ | <u>23.16</u>    | 0.63            | 3.21              |
| UniTok* [25]                | $\checkmark$ | $16 \times 16$ | <b>25.32</b>    | <b>0.77</b>     | <b>0.37</b>       |
| ViQ                         | $\checkmark$ | $16 \times 16$ | 22.73           | <u>0.66</u>     | <u>0.62</u>       |

aggressive quality setting (e.g.,  $Q \approx 0.08$ ) that noticeably degrades image quality, whereas ViQ preserves substantially better reconstruction at an equivalent ratio. We therefore compare ViQ against JPEG at a comparable bitrate, and report the compression ratios together with reconstructed visual samples in Figure 4.

#### 4.4 Image Reconstruction Experiments

*Setups.* For image reconstruction tasks, we train the visual decoder with the fixed pre-trained ViQ model. The overall regularization loss for VAE is defined as:

$$\mathcal{L}_{\text{VAE}} = \mathcal{L}_{\text{KL}} + \mathcal{L}_{\text{MSE}} + \mathcal{L}_{\text{LPIPS}} + \mathcal{L}_{\text{GAN}} + \lambda_{\text{REPA}} \cdot \mathcal{L}_{\text{REPA}} \quad (8)$$

where  $\mathcal{L}_{\text{KL}}$  denotes the Kullback-Leibler divergence,  $\mathcal{L}_{\text{MSE}}$  represents the Mean Squared Error,  $\mathcal{L}_{\text{LPIPS}}$  corresponds to the Learned Perceptual Image Patch Similarity for pixel reconstruction, and  $\mathcal{L}_{\text{GAN}}$  is the adversarial loss from the Generative Adversarial Network [12]. To enhance semantic alignment, we leverage

DINOv2 [30] as an external feature extractor and apply an alignment loss  $\mathcal{L}_{\text{REPA}}$  to the ViQ preprocess layer preceding the decoder head within our ViQ architecture. The alignment loss coefficient is empirically set to  $\lambda_{\text{REPA}} = 1.5$  for the DINOv2 head. For model optimization, we employ the AdamW optimizer [15, 23] with a constant learning rate of  $6 \times 10^{-4}$  and a global batch size of 4096. Training stability is ensured through gradient clipping and the application of an Exponential Moving Average (EMA) to the generative model parameters. All models are trained for 50,000 steps with a linear learning rate warmup over the initial 5,000 steps. We utilize mixed-precision training with the BF16 format and unfreeze the VAE preprocess layer during training.

*Results.* We comprehensively evaluate the reconstruction quality of various tokenizers on the  $256 \times 256$  ImageNet-1K validation set, as summarized in Table 2. Among discrete tokenizers, ViQ delivers highly competitive reconstruction quality, attaining an SSIM of 0.66 and an rFID of 0.62 (second only to UniTok), while remaining comparable to QLIP-B in PSNR (22.73 vs. 23.16). We note that UniTok reports stronger raw reconstruction metrics (e.g., PSNR 25.32), which we attribute to its direct reconstruction objective combined with contrastive supervision. However, as shown in Table 1, this type of supervision substantially degrades vision-language alignment and thus leads to sub-optimal multi-modal understanding. By contrast, our core objective with ViQ is to deliver a more balanced tokenizer that maintains highly competitive low-level reconstruction while excelling in high-level semantic understanding. This balance indicates that a well-optimized pixel-based decoder remains a powerful and efficient choice for building a discrete visual tokenizer that serves both generative and discriminative purposes. We also acknowledge that, as a discrete tokenizer, ViQ inevitably compromises some fine-grained reconstruction detail relative to continuous tokenizers, as the continuous visual feature space is drastically compressed; nevertheless, it considerably narrows this gap and offers a favorable trade-off between reconstruction fidelity and semantic understanding.

#### 4.5 Ablation Studies

**Proximal Representations** In Tables 3a and 3b, we show that initializing the model with intermediate proximal representations can substantially improve the final performance. In Table 3a, we can see that optimizing a continuous feature space into a discrete one leads to a severe performance drop. We further investigate the effect of introducing a bottleneck, a bottleneck with L2 normalization, and a bottleneck with  $L_\infty$  normalization. The results highlight two key observations: first, a gradually regularized latent space is more suitable for quantization; second, a more appropriate regularization method for quantization (i.e.,  $L_\infty$ ) provides additional benefits. Furthermore, Table 3b shows that, when optimized properly, the bottleneck does not necessarily degrade the final performance.

**Quantization Designs** We compared the performance of commonly used methods such as SimVQ [53] and FSQ [29] under different codebook sizes, as shown

| Process                                                  | avg.(2-2) | Width | avg.(2-1) |
|----------------------------------------------------------|-----------|-------|-----------|
| Continuous $\rightarrow$ SimVQ                           | 60.9      | 32    | 68.4      |
| Continuous $\rightarrow$ BN + $l_2 \rightarrow$ SimVQ    | 66.6      | 128   | 69.1      |
| Continuous $\rightarrow$ BN + $l_2 \rightarrow$ FSQ      | 67.9      | 512   | 68.8      |
| Continuous $\rightarrow$ BN + $l_\infty \rightarrow$ FSQ | 68.7      | 1536  | 69.3      |

(a) **Proximal Representations.** Gradually regularizing the latent space from continuous to quantized helps.

(b) **Bottleneck size.** The bottleneck width can be made significantly narrower than the 1536-dimensional one in SigLIP2-g.

| Size             | avg.(2-2) | Design                           | avg.(2-2) |
|------------------|-----------|----------------------------------|-----------|
| FSQ + 64,000     | 68.7      | no Position information injected | 65.3      |
| FSQ + 128,000    | 68.3      | RoPE with Attention              | 68.7      |
| SimVQ + $2^{15}$ | 66.5      | Learnable Pos Emb                | 65.7      |
| SimVQ + $2^{17}$ | 65.6      |                                  |           |

(c) **VQ and Codebook size.** The VQ codebook saturates at a size of  $2^{17}$ .

(d) **Position encoding for Quantization.** position information is important.

| Case | Self-Distill | Text | Recon | avg.(2-2) | Loss Type       | Time Cost | avg.(2-2) |
|------|--------------|------|-------|-----------|-----------------|-----------|-----------|
| A    | ✗            | ✓    | ✗     | 61.3      | none            | 1x        | 66.8      |
| B    | ✓            | ✓    | ✗     | 66.8      | MSE + LPIPS     | 2.3x      | 67.0      |
| C    | ✓            | ✓    | ✓     | 68.7      | DiT (unfrozen)  | 4x        | 65.8      |
|      |              |      |       |           | DiT (frozen)    | 4x        | 67.6      |
|      |              |      |       |           | Vae Latent Loss | 1.3x      | 68.7      |

(e) **Loss Combination.** three type of loss.

(f) **Reconstruction Loss.** Vae Latent loss is efficient.

**Table 3: ViQ ablation experiments.** We conduct a thorough ablation study of ViQ, examining proximal representations, architecture design, and loss combinations. We report the average metrics on MMStar, MMMU, OCRBench, among a total of eight benchmarks, using a fixed SFT training setup on Qwen2.5-3B. (2-1) and (2-2) refer to the training stage of the test model. Default settings are marked in red.

in Table 3c. We conduct experiments on different position encoding methods for ViQ, summarized in Table 3d. We observe that FSQ [29], a non-optimized vector quantization method, outperforms SimVQ [53], which requires learning a codebook. We also experimented with LFQ [47], vanilla VQ [11], and IBQ [33], and found consistent conclusions: in our setting, VQ methods without a learnable vocabulary tend to perform better. In addition, we find that a vocabulary size around 60,000 is good. Increasing the codebook size will reduce the utilization rate of learnable codebooks, which in turn degrades performance, whereas this issue is much less pronounced for non-learnable VQ methods. Table 3d further shows that injecting positional encoding significantly enhances the representational ability of VQ. We conduct experiments on both learnable positional embeddings and RoPE-2D, and observe that RoPE leads to a substantial improvement, whereas learnable positional embeddings offer limited gains. This is because learnable positional embeddings increase the optimization difficulty of the VQ module.

**Loss Combination.** We further study the loss composition used in our training pipeline. As shown in Table 3e. When keep only the text loss, it leads to a substantial performance drop (Case A). In Case B, add self distillation loss with significant enhance the performace of model. While we can see some unsatisfactory performance on tasks such as OCR and Chart. By introducing a reconstruction loss (Case C), the performance notably improves, with most gains coming from detail-intensive tasks such as OCR and Chart, while the gains for captioning and VQA are relatively limited. Table 3f studies different formulations of the reconstruction loss. We adopt a more efficient and simple approach: a VAE latent reconstruction loss, where the model predicts the latent representation of a pretrained VAE network. This achieves effective optimization while maintaining training efficiency. We further analyze why the alternative formulations are less effective. For the DiT-based objective, jointly training an *unfrozen* DiT module lets the generative branch absorb part of the representational capacity of the encoder, which degrades the discriminative features and even drops below the no-reconstruction baseline (65.8); freezing the DiT alleviates this and restores the score to 67.6, yet it still trails the VAE latent loss due to the lack of fine-tuning flexibility. For pixel-level objectives, MSE and LPIPS mainly enforce rigid pixel-wise fitting, which tends to compete with and erode the high-level semantic features that are crucial for multi-modal understanding. By contrast, supervising in the latent space of a pretrained VAE strikes a better balance: it guides low-level detail reconstruction while preserving semantic integrity, and does so at a much lower computational cost ( $1.3\times$ ).

## 5 Conclusion

In this work, we introduced ViQ, a quantized multimodal encoder that unifies visual and language representations at native resolution. We designed a two-stage training pipeline to reduce the information losses of learning discrete visual representations with text-aligned pretraining and feature discretization. Through carefully designed architectures and progressive training techniques that minimize information loss, ViQ achieves competitive performance in multimodal tasks compared to both existing continuous and discrete visual encoders. We believe our work not only enhances the quality and efficiency of visual representations in multimodal learning, but also offers a viable pathway toward unifying vision and language within a shared discrete framework.

## Acknowledgement

This work was supported in part by the National Key Research and Development Program of China under Grant 2023YFB280690, and in part by the National Natural Science Foundation of China under Grant 62321005, Grant 62336004, and Grant 62125603.

## References

1. Agarwal, N., Ali, A., Bala, M., Balaji, Y., Barker, E., Cai, T., Chattopadhyay, P., Chen, Y., Cui, Y., Ding, Y., et al.: Cosmos world foundation model platform for physical ai. arXiv preprint arXiv:2501.03575 (2025)
2. Chen, L., Li, J., Dong, X., Zhang, P., Zang, Y., Chen, Z., Duan, H., Wang, J., Qiao, Y., Lin, D., et al.: Are we on the right way for evaluating large vision-language models? arXiv preprint arXiv:2403.20330 (2024)
3. Chen, Z., Wang, W., Cao, Y., Liu, Y., Gao, Z., Cui, E., Zhu, J., Ye, S., Tian, H., Liu, Z., et al.: Expanding performance boundaries of open-source multimodal models with model, data, and test-time scaling, 2025. URL <https://arxiv.org/abs/2412.05271>
4. Chen, Z., Wu, J., Wang, W., Su, W., Chen, G., Xing, S., Zhong, M., Zhang, Q., Zhu, X., Lu, L., et al.: Internvl: Scaling up vision foundation models and aligning for generic visual-linguistic tasks. In: CVPR. pp. 24185–24198 (2024)
5. Cheng, X., Zhang, W., Zhang, S., Yang, J., Guan, X., Wu, X., Li, X., Zhang, G., Liu, J., Mai, Y., et al.: Simplevqa: Multimodal factuality evaluation for multimodal large language models. In: Proceedings of the IEEE/CVF International Conference on Computer Vision. pp. 4637–4646 (2025)
6. Chiu, C.C., Qin, J., Zhang, Y., Yu, J., Wu, Y.: Self-supervised learning with random-projection quantizer for speech recognition. In: International Conference on Machine Learning. pp. 3915–3924. PMLR (2022)
7. Dehghani, M., Mustafa, B., Djolonga, J., Heek, J., Minderer, M., Caron, M., Steiner, A., Puigcerver, J., Geirhos, R., Alabdulmohsin, I.M., et al.: Patch n’pack: Navit, a vision transformer for any aspect ratio and resolution. NeurIPS **36** (2024)
8. 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)
9. Fang, A., Jose, A.M., Jain, A., Schmidt, L., Toshev, A., Shankar, V.: Data filtering networks. arXiv preprint arXiv:2309.17425 (2023)
10. Fini, E., Shukor, M., Li, X., Dufter, P., Klein, M., Haldimann, D., Aitharaju, S., da Costa, V.G.T., Béthune, L., Gan, Z., et al.: Multimodal autoregressive pre-training of large vision encoders. In: Proceedings of the Computer Vision and Pattern Recognition Conference. pp. 9641–9654 (2025)
11. Gersho, A., Gray, R.M.: Vector quantization and signal compression, vol. 159. Springer Science & Business Media (2012)
12. Goodfellow, I., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A., Bengio, Y.: Generative adversarial networks. Communications of the ACM **63**(11), 139–144 (2020)
13. Guo, D., Wu, F., Zhu, F., Leng, F., Shi, G., Chen, H., Fan, H., Wang, J., Jiang, J., Wang, J., et al.: Seed1. 5-v1 technical report. arXiv preprint arXiv:2505.07062 (2025)
14. Kembhavi, A., Salvato, M., Kolve, E., Seo, M., Hajishirzi, H., Farhadi, A.: A diagram is worth a dozen images. In: ECCV. pp. 235–251. Springer (2016)
15. Kingma, D.P., Ba, J.: Adam: A method for stochastic optimization (2017), <https://arxiv.org/abs/1412.6980>
16. Laurençon, H., Tronchon, L., Cord, M., Sanh, V.: What matters when building vision-language models? arXiv preprint arXiv:2405.02246 (2024)
17. Li, B., Zhang, P., Zhang, K., Pu, F., Du, X., Dong, Y., Liu, H., Zhang, Y., Zhang, G., Li, C., et al.: Lmms-eval: Accelerating the development of large multimodal models (2024)



18. Li, B., Zhang, Y., Guo, D., Zhang, R., Li, F., Zhang, H., Zhang, K., Li, Y., Liu, Z., Li, C.: Llava-onevision: Easy visual task transfer. arXiv preprint arXiv:2408.03326 (2024)
19. Liu, H., Li, C., Li, Y., Li, B., Zhang, Y., Shen, S., Lee, Y.J.: Llava-next: Improved reasoning, ocr, and world knowledge (2024)
20. Liu, H., Li, C., Wu, Q., Lee, Y.J.: Visual instruction tuning. *NeurIPS* **36** (2024)
21. Liu, Y., Li, Z., Yang, B., Li, C., Yin, X., Liu, C.I., Jin, L., Bai, X.: On the hidden mystery of ocr in large multimodal models. arXiv preprint arXiv:2305.07895 (2023)
22. Liu, Z., Dong, Y., Liu, Z., Hu, W., Lu, J., Rao, Y.: Oryx mllm: On-demand spatial-temporal understanding at arbitrary resolution. arXiv preprint arXiv:2409.12961 (2024)
23. Loshchilov, I., Hutter, F.: Decoupled weight decay regularization. arXiv preprint arXiv:1711.05101 (2017)
24. Luo, Z., Shi, F., Ge, Y., Yang, Y., Wang, L., Shan, Y.: Open-magvit2: An open-source project toward democratizing auto-regressive visual generation. arXiv preprint arXiv:2409.04410 (2024)
25. Ma, C., Jiang, Y., Wu, J., Yang, J., Yu, X., Yuan, Z., Peng, B., Qi, X.: Uniotok: A unified tokenizer for visual generation and understanding. arXiv preprint arXiv:2502.20321 (2025)
26. Masry, A., Long, D.X., Tan, J.Q., Joty, S., Hoque, E.: Chartqa: A benchmark for question answering about charts with visual and logical reasoning. arXiv preprint arXiv:2203.10244 (2022)
27. Mathew, M., Bagal, V., Tito, R., Karatzas, D., Valveny, E., Jawahar, C.: Infographicvqa. In: *Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision*. pp. 1697–1706 (2022)
28. Mathew, M., Karatzas, D., Jawahar, C.: Docvqa: A dataset for vqa on document images. In: *Proceedings of the IEEE/CVF winter conference on applications of computer vision*. pp. 2200–2209 (2021)
29. Mentzer, F., Minnen, D., Agustsson, E., Tschannen, M.: Finite scalar quantization: Vq-vae made simple. In: *The Twelfth International Conference on Learning Representations*
30. Quab, M., Darcet, T., Moutakanni, T., Vo, H., Szafraniec, M., Khalidov, V., Fernandez, P., Haziza, D., Massa, F., El-Nouby, A., et al.: Dinov2: Learning robust visual features without supervision. arXiv preprint arXiv:2304.07193 (2023)
31. QwenTeam: Qwen2-vl: To see the world more clearly. Wwen Blog (2024), <https://qwenlm.github.io/blog/qwen2-vl/>
32. Radford, A., Kim, J.W., Hallacy, C., Ramesh, A., Goh, G., Agarwal, S., Sastry, G., Askell, A., Mishkin, P., Clark, J., et al.: Learning transferable visual models from natural language supervision. In: *ICML*. pp. 8748–8763. PMLR (2021)
33. Shi, F., Luo, Z., Ge, Y., Yang, Y., Shan, Y., Wang, L.: Scalable image tokenization with index backpropagation quantization. In: *Proceedings of the IEEE/CVF International Conference on Computer Vision*. pp. 16037–16046 (2025)
34. Singh, A., Natarajan, V., Shah, M., Jiang, Y., Chen, X., Batra, D., Parikh, D., Rohrbach, M.: Towards vqa models that can read. In: *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*. pp. 8317–8326 (2019)
35. Su, J., Ahmed, M., Lu, Y., Pan, S., Bo, W., Liu, Y.: Roformer: Enhanced transformer with rotary position embedding. *Neurocomputing* **568**, 127063 (2024)
36. 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)

37. Team, K., Du, A., Yin, B., Xing, B., Qu, B., Wang, B., Chen, C., Zhang, C., Du, C., Wei, C., et al.: Kimi-vl technical report. arXiv preprint arXiv:2504.07491 (2025)
38. Team, M.L.: Longcat-next: Lexicalizing modalities as discrete tokens (2026), <https://arxiv.org/abs/2603.27538>
39. Team, Q.: Qwen2.5-vl (January 2025), <https://qwenlm.github.io/blog/qwen2.5-vl/>
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. Wan, T., Wang, A., Ai, B., Wen, B., Mao, C., Xie, C.W., Chen, D., Yu, F., Zhao, H., Yang, J., et al.: Wan: Open and advanced large-scale video generative models. arXiv preprint arXiv:2503.20314 (2025)
43. Wu, C., Li, J., Zhou, J., Lin, J., Gao, K., Yan, K., Yin, S.m., Bai, S., Xu, X., Chen, Y., et al.: Qwen-image technical report. arXiv preprint arXiv:2508.02324 (2025)
44. Xie, J., Mao, W., Bai, Z., Zhang, D.J., Wang, W., Lin, K.Q., Gu, Y., Chen, Z., Yang, Z., Shou, M.Z.: Show-o: One single transformer to unify multimodal understanding and generation. arXiv preprint arXiv:2408.12528 (2024)
45. Xie, R., Du, C., Song, P., Liu, C.: Muse-vl: Modeling unified vlm through semantic discrete encoding. In: *Proceedings of the IEEE/CVF International Conference on Computer Vision*. pp. 24135–24146 (2025)
46. Yin, W., Ye, Y., Shu, F., Liao, Y., Kang, Z., Dong, H., Yu, H., Yang, D., Wang, J., Wang, H., et al.: Sail-vl2 technical report. arXiv preprint arXiv:2509.14033 (2025)
47. Yu, L., Lezama, J., Gundavarapu, N.B., Versari, L., Sohn, K., Minnen, D., Cheng, Y., Gupta, A., Gu, X., Hauptmann, A.G., et al.: Language model beats diffusion-tokenizer is key to visual generation. In: *The Twelfth International Conference on Learning Representations*
48. Yue, X., Ni, Y., Zhang, K., Zheng, T., Liu, R., Zhang, G., Stevens, S., Jiang, D., Ren, W., Sun, Y., et al.: Mmmu: A massive multi-discipline multimodal understanding and reasoning benchmark for expert agi. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 9556–9567 (2024)
49. 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)
50. Zhao, Y., Xiong, Y., Kraehenbuehl, P.: Image and video tokenization with binary spherical quantization. In: *The Thirteenth International Conference on Learning Representations*
51. Zhao, Y., Xue, F., Reed, S., Fan, L., Zhu, Y., Kautz, J., Yu, Z., Krähenbühl, P., Huang, D.A.: Qlip: Text-aligned visual tokenization unifies auto-regressive multimodal understanding and generation. *CoRR* (2025)
52. Zhu, J., Wang, W., Chen, Z., Liu, Z., Ye, S., Gu, L., Tian, H., Duan, Y., Su, W., Shao, J., et al.: Internvl3: Exploring advanced training and test-time recipes for open-source multimodal models. arXiv preprint arXiv:2504.10479 (2025)
53. Zhu, Y., Li, B., Xin, Y., Xia, Z., Xu, L.: Addressing representation collapse in vector quantized models with one linear layer. In: *Proceedings of the IEEE/CVF International Conference on Computer Vision*. pp. 22968–22977 (2025)