

From Specific-MLLMs to Omni-MLLMs: A Survey on MLLMs Aligned with Multi-modalities

Anonymous ACL submission

Abstract

To tackle complex tasks in real-world scenarios, more researchers are focusing on Omni-MLLMs, which aim to achieve omni-modal understanding and generation. Beyond the constraints of any specific non-linguistic modality, Omni-MLLMs map various non-linguistic modalities into the embedding space of LLMs and enable the interaction and understanding of arbitrary combinations of modalities within a single model. In this paper, we systematically investigate relevant research and provide a comprehensive survey of Omni-MLLMs. Specifically, we first explain the four core components of Omni-MLLMs for unified multi-modal modeling with a meticulous taxonomy that offers novel perspectives. Then, we introduce the effective integration achieved through two-stage training and discuss the corresponding datasets as well as evaluation. Furthermore, we summarize the main challenges of current Omni-MLLMs and outline future directions. We hope this paper serves as an introduction for beginners and promotes the advancement of related research. Resources will be made public.

1 Introduction

The remarkable performance of continuously evolving Multi-modal Large Language Models (MLLMs) has pointed to a possible direction for achieving general artificial intelligence (Bubeck et al., 2023; OpenAI, 2023b). MLLMs extend Large Language Models (LLMs) by integrating them with pre-trained models tailored to specific modalities, such as Vision-MLLMs (Liu et al., 2023c; Wang et al., 2024b; Sun et al., 2024c), Audio-MLLMs (Zhang et al., 2023a; Chu et al., 2023), and 3D-MLLMs (Xu et al., 2024b). However, these modality-specific MLLMs (Specific-MLLMs) are insufficient to tackle complex tasks in real-world scenarios that simultaneously involve multiple modalities. Therefore, efforts are being made to expand the range of modalities for

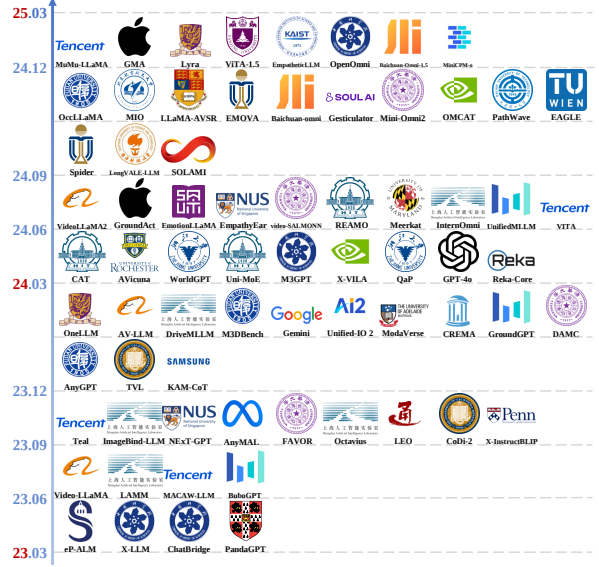


Figure 1: The timeline of representative Omni-MLLMs.

understanding and generation, giving rise to the omni-modality MLLMs (Omni-MLLMs).

By integrating multiple pre-trained models of more non-linguistic modalities (Radford et al., 2021, 2023; Xue et al., 2024b; Rombach et al., 2022; Liu et al., 2023b), Omni-MLLMs expand the modalities for understanding and even generation based on Specific-MLLMs. Omni-MLLMs leverage the emergent capabilities of LLMs to treat various non-linguistic modalities as different *foreign languages*, enabling the interaction and understanding of information across different modalities within a unified space (Chen et al., 2023a; Panagopoulou et al., 2024). Compared to Specific-MLLMs, Omni-MLLMs can perform multiple uni-modal¹ understanding and generation tasks, as well as cross-modal tasks across two or more non-linguistic modalities, allowing a single model to

¹“Uni” and “Cross” refer to the number of non-linguistic modalities involved in the interaction, in contrast to “multi-modal reasoning,” traditionally reserved for vision-language tasks (Panagopoulou et al., 2024).

handle arbitrary combinations of modalities.

A review of the development of Omni-MLLMs reveals that it has been continuously expanding in three directions. On the one hand, the types of modalities processed by Omni-MLLM have been continuously increasing, from X-LLMs that handle vision and audio to X-InstructBLIP (Panagopoulou et al., 2024) which adds 3D modality capabilities, PandaGPT (Su et al., 2023) that incorporates IMU modality, and finally One-LLM (Han et al., 2024a), which processes eight different modalities simultaneously. On the other hand, the ability to interact across modalities of Omni-MLLMs has also expanded, from the joint 3D-Image and Audio-Image cross-modal reasoning capability in ImageBind-LLM (Han et al., 2023) to the cross-modal generation capability of CoDi-2 that leverages interleaved audio and image contexts to generate both audio and images (Tang et al., 2024b). The Omni-MLLM is thus trending towards an “Any-to-Any” model. Besides, the application scenarios of Omni-MLLMs have been broadened, encompassing real-time multimodal speech interaction like Mini-Omni2 and Lyra (Xie and Wu, 2024; Zhong et al., 2024), world simulation like WordGPT (Ge et al., 2024b), multi-sensor autonomous driving like DriveMLM (Wang et al., 2023b), etc. In addition to the open-source models, there are also some closed-source Omni-MLLMs such as GPT-4o (OpenAI), Gemini (Reid et al., 2024), and Reka (Ormazabal et al., 2024). The timeline of Omni-MLLMs is shown in Figure 1. Despite the emergence of numerous Omni-MLLMs, there is still a lack of systematic evaluation and analysis.

To fill the gap, we propose this work to conduct a comprehensive and detailed analysis of Omni-MLLMs. We first review the architecture of Omni-MLLMs in four parts (§2). Next, we summarize how Omni-MLLMs expand across multiple modalities through the two-stage training process (§3); then present the training data construction and performance evaluation (§4). Furthermore, we highlight some key challenges and future directions (§5). Finally, we provide a brief summary (§6) and discuss related surveys in the Appendix A.

Our contributions can be summarized as follows:

- (1) **Comprehensive Survey**: This is the first comprehensive survey dedicated for Omni-MLLMs;
- (2) **Meticulous taxonomy**: We introduce a meticulous taxonomy (shown in Figure 2);
- (3) **Challenges and Future**: We outline the challenges of Omni-MLLMs and shed light on future research.

2 Omni-MLLM Architecture

As the extension of Specific-MLLMs, Omni-MLLMs inherit the architecture of *encoding, alignment, interaction, and generation* and broaden the types of non-linguistic modalities involved. This section introduces the implementation methods and functions of the four components in Omni-MLLM: Multi-modalities Encoding (§2.1), Multi-modalities Alignment (§2.2), Multi-modalities Interaction (§2.3), and Multi-modalities Generation (§2.4). More details about the architecture of Omni-MLLMs are shown in Appendix B.

2.1 Multi-modalities Encoding

Based on the encoding feature spaces of multiple modalities, we categorize the Omni-MLLM encoding methods into three types: 1) continuous encoding, 2) discrete encoding, and 3) hybrid encoding.

2.1.1 Continuous Encoding

Continuous encoding refers to encoding the modality into the continuous feature space. Omni-MLLMs that adopt continuous encoding, such as X-LLM (Chen et al., 2023a) and ChatBridge (Zhao et al., 2023b), often integrate multiple pre-trained uni-modality encoders. These modality-specific encoders encode different modalities $\mathbf{X}$ into distinct feature spaces $\mathbb{R}_x$ as $\mathbf{F}_x$, formulated as:

$$\mathbf{F}_x = \text{SpecificEncoder}(\mathbf{X}), \mathbf{F}_x \in \mathbb{R}_x \quad (1)$$

where SpecificEncoder refers to different modality-specific encoders used in Omni-MLLMs, such as InternVit (Chen et al., 2023g) for encoding visual modality, Whisper (Radford et al., 2023) for encoding auditory modality, ULIP-2 (Xue et al., 2024b) for encoding 3D modality, IMU2CLIP (Moon et al., 2022) for encoding IMU modality, etc.

Besides using heterogeneous encoders for continuous encoding, some Omni-MLLMs (Han et al., 2024a, 2023; Su et al., 2023; Fu et al., 2024c) employ pre-aligned encoders for multiple modalities, encoding different modalities $\mathbf{X}$ into the same feature space $\mathbb{R}_{uni}$, as shown in Equation 2.

$$\mathbf{F}_x = \text{PreAlignEncoder}(\mathbf{X}), \mathbf{F}_x \in \mathbb{R}_{uni} \quad (2)$$

where PreAlignEncoder refer to encoders that uniformly encode multiple modalities, such as LanguageBind (Zhu et al., 2024a) which uses text as a bridge to align different modalities, and ImageBind (Girdhar et al., 2023) which uses images as a bridge to align different modalities.

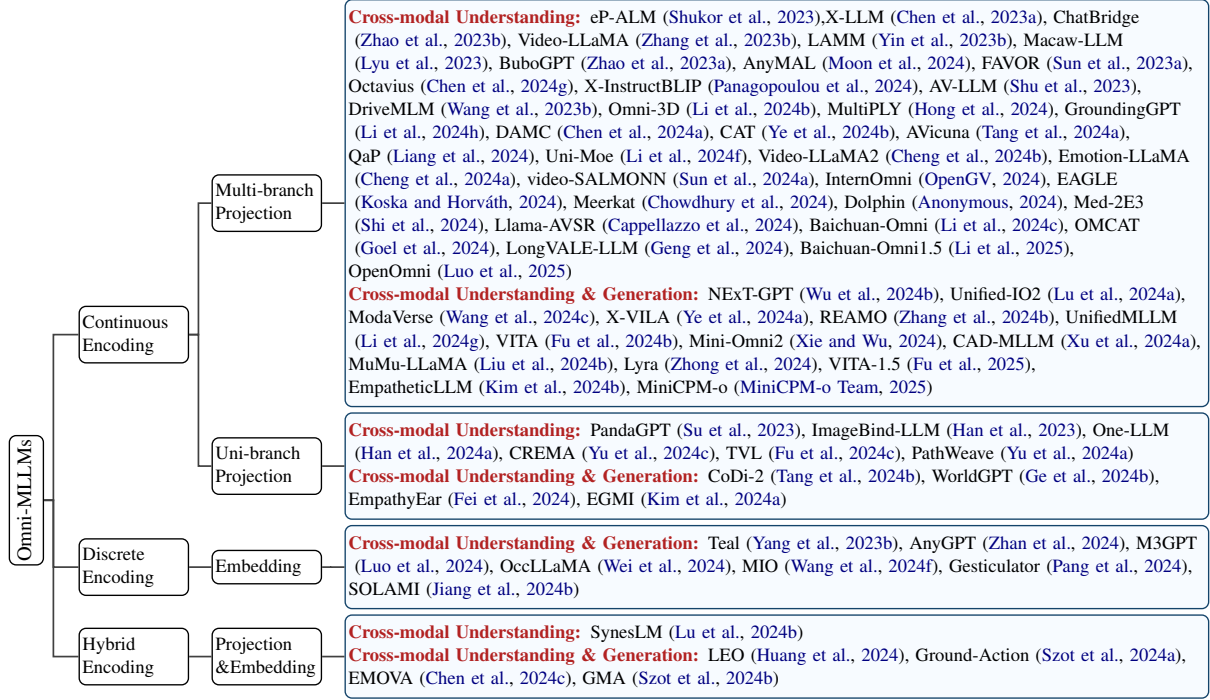


Figure 2: Taxonomy for Omni-MLLMs based on their encoding and alignment methods.

2.1.2 Discrete Encoding

To better facilitate the seamless integration and generation of new non-linguistic modalities, some Omni-MLLMs, such as AnyGPT (Zhan et al., 2024) and Teal (Yang et al., 2023b), adopt a discrete encoding approach. This method encodes different raw modalities $\mathbf{X}$ into the same discrete token space $\mathbb{V}_{uni}$ as $\mathbf{T}_x$, formulated as follows:

$$\mathbf{T}_x = \text{SpecificTokenizer}(\mathbf{X}), \quad \mathbf{T}_x \in \mathbb{V}_{uni} \quad (3)$$

where *SpecificTokenizer* refers to different modality-specific tokenizers used in Omni-MLLMs, including the SEED tokenizer (Ge et al., 2024a) based on Vector Quantized Tokenization (VQ), the SpeechTokenizer (Zhang et al., 2023c) based on Residual Vector Quantized Tokenization (RVQ), the AudioTokenizer of Teal (Yang et al., 2023b) based on k-means clustering, and so on.

2.1.3 Hybrid Encoding

Although discrete encoding facilitates the unified processing of different non-linguistic modalities and text compared to continuous encoding, discrete modality tokens often struggle to capture the detailed information inherent in raw continuous modalities (Chen et al., 2024c; Xie and Wu, 2024). Therefore, some Omni-MLLMs combine both encoding approaches instead of a fully discretized manner, choosing different encoding methods for

different modalities. For instance, EMOVA (Chen et al., 2024c) uses the discrete S2U tokenizer to encode auditory modalities while employing the continuous encoder InternViT for visual modalities to retain more vision semantic information. Similarly, GroundAction (Szot et al., 2024a) encodes visual modalities using the CLIP ViT and action modalities with its trained action tokenizer.

2.2 Multi-modalities Alignment

Omni-MLLMs align the encoded features of various non-linguistic modalities with the embedding space of LLMs. The multi-modality alignment can be categorized into two approaches: 1) projection alignment and 2) embedding alignment.

2.2.1 Projection Alignment

The continuous encoding Omni-MLLMs insert adapters, referred to as *projectors*, between the encoders and the LLMs. These projectors map the continuously encoded modality features $\mathbf{F}_x$ into the text embedding space as $\mathbf{F}_p$. As discussed in Section 2.1.1, $\mathbf{F}_x$ may either reside in distinct feature spaces $\mathbb{R}_x$ or share the same feature space $\mathbb{R}_{uni}$. For the former, multiple projectors are typically employed to align the $\mathbf{F}_x$ of each modality into $\mathbb{R}_t$ as $\mathbf{F}_p$ independently, addressing dimensional mismatch and feature misalignment across modalities (Ye et al., 2024a; Lyu et al., 2023; Moon et al.,

2024), formulated as follows:

$$\mathbf{F}_p = \text{SpecificProjector}(\mathbf{F}_x), \mathbf{F}_p \in \mathbb{R}_t \quad (4)$$

where `SpecificProjector` refers to the modality-specific projector corresponding to different modalities, called *multi-branch projection*.

For the latter case, besides the multi-branch approach, Omni-MLLMs like PandaGPT (Su et al., 2023) and WorldGPT (Ge et al., 2024b) adopt a shared projector to achieve unified alignment across modalities to reduce the parameters of multiple projectors, as shown in Equation 5.

$$\mathbf{F}_p = \text{UnifiedProjector}(\mathbf{F}_x), \mathbf{F}_p \in \mathbb{R}_t \quad (5)$$

where `UnifiedProjector` refers to the unified projector used to align multiple modalities, a design known as the *uni-branch projection*. A comparison of the two approaches is illustrated in Figure 3.

In terms of the *specific implementation* of the projector, the most straightforward approach is to use a multi-layer perceptron (MLP) or a single linear layer (Wu et al., 2024b; Cheng et al., 2024b; OpenGV, 2024). Alternatively, attention mechanisms can be employed to compress the encoded information of non-linguistic modalities. This includes cross-attention-based methods like Q-Former (Panagopoulou et al., 2024; Chen et al., 2023a) and Perceiver (Zhao et al., 2023b; Liang et al., 2024), as well as self-attention-based methods such as UPM in OneLLM (Han et al., 2024a). Additionally, BaiChuan-Omni (Li et al., 2024c) and EMOVA (Chen et al., 2024c) incorporate CNNs to compress the projected features, thereby achieving locality preservation (Cha et al., 2024).

It is also worth noting that in multi-branch Omni-MLLMs, different branches may utilize distinct implementations to better accommodate the unique characteristics of each modality (Li et al., 2024h). For example, Uni-MoE (Li et al., 2024f) uses a linear projection for the visual modality and a Q-Former for the auditory modality. Meanwhile, uni-branch Omni-MLLMs, when using an attention-based projector, typically design multiple modality-specific learnable vectors to extract key information from various non-linguistic modalities (Yu et al., 2024a; Han et al., 2024a; Yu et al., 2024c).

2.2.2 Embedding Alignment

As for discrete encoding Omni-MLLMs, the features of non-linguistic modalities are represented as quantized codes, which reside in the same discrete space $\mathbb{V}_{uni}$ as text tokens. Therefore, new

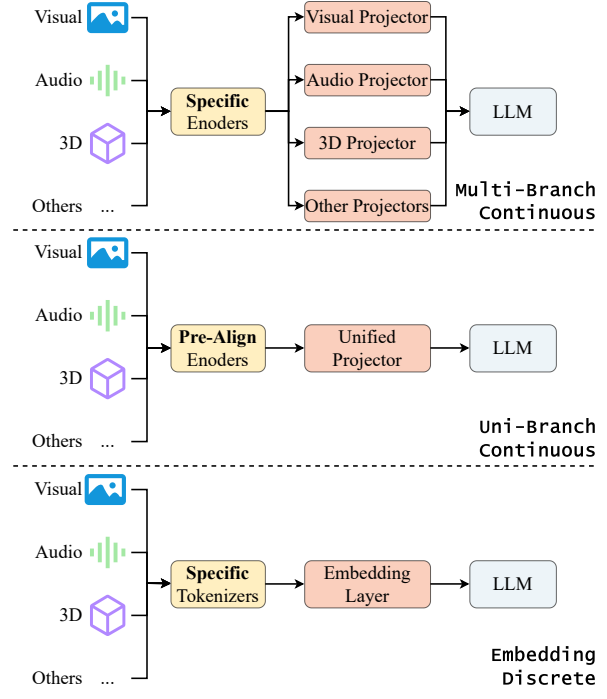


Figure 3: The three combinations of encoding and alignment in Omni-MLLM are based on different encoding spaces and alignment structures.

modality-specific discrete tokens $\mathbf{T}_x$ are embedded into the continuous feature space $\mathbb{R}_t$ by modifying the vocabulary of LLMs and the corresponding embeddings layer, as shown in Equation 6.

$$\mathbf{F}_p = \text{Embedding}(\mathbf{T}_x), \mathbf{F}_p \in \mathbb{R}_t \quad (6)$$

where `Embedding` refers to the unified embedding layer corresponding to different modalities, which is typically achieved by adding discrete codebooks from various modalities to the vocabulary and expanding the embedding layer of LLMs (Zhan et al., 2024; Yang et al., 2023b; Wei et al., 2024). For instance, AnyGPT extends the vocabulary of the LLaMA-2 by incorporating 17,408 codes across three modalities—image, speech, and music (Zhan et al., 2024). Besides, some works like Ground-Action (Szot et al., 2024a) and LEO (Huang et al., 2024) overwrite infrequently used tokens in the original vocabulary for alignment, as they extend a smaller set of modality-specific discrete tokens.

Additionally, for hybrid encoding models, alignment is achieved by simultaneously employing both the projection method and the embedding method (Chen et al., 2024c; Szot et al., 2024b).

2.3 Multi-modalities Interaction

Omni-MLLMs utilize transformer-based LLMs to facilitate information interaction between different

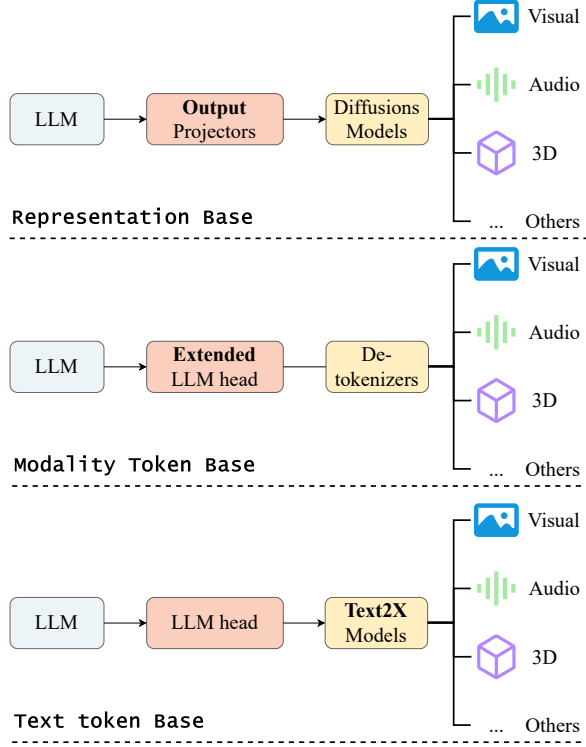


Figure 4: Three generation methods in Omni-MLLMs are implemented based on different output spaces of LLM and the corresponding generative models.

modalities within a unified feature space $\mathbb{R}_t$. Commonly used LLMs include the LLaMA series (Touvron et al., 2023), the Qwen series (Bai et al., 2023), and others (Cai et al., 2024; Zeng et al., 2024).

For interaction, most Omni-MLLMs (Chen et al., 2023a; Ye et al., 2024a; OpenGV, 2024; Li et al., 2025) concatenate aligned non-linguistic modality features $\mathbf{F}_p$ with textual features $\mathbf{F}_t$ at the input level, enabling interaction in a progressive and layer-by-layer manner. Meanwhile, some works, such as ImageBind-LLM (Han et al., 2023) and TVL-LLaMA (Fu et al., 2024c), insert $\mathbf{F}_p$ into specific layers or all layers of the LLMs to mitigate the loss of modality information (Shukor et al., 2023).

In terms of the number of modalities involved in interactions, compared to Specific-MLLMs that are limited to dual-modal interactions between a single non-linguistic modality and text (Liu et al., 2023c; Xu et al., 2024b), Omni-MLLMs not only support multiple dual-modal interactions but also enable omni-multimodal interactions involving more than two non-linguistic modalities (Zhao et al., 2023b; Wang et al., 2024f). For example, X-InstructBLIP (Panagopoulou et al., 2024) enables dual-modal interactions such as vision-text, audio-text, and 3D-text, as well as omni-modal inter-

actions like vision-audio-text and 3D-vision-text, showcasing the ability of Omni-MLLMs to handle arbitrary combinations of modalities.

2.4 Multi-modalities Generation

Omni-MLLMs can output text while also generating non-linguistic modalities by integrating different generation models. As shown in Figure 4, we categorize multi-modalities generation into three types: text-based generation, representation-based generation, and modality-token-based generation.

Text-based This approach directly utilizes the discrete text output from the LLM to invoke Text-to-X generation models (Liu et al., 2023b; Luo et al., 2023c; Brooks et al., 2023) based on the content of the text. For example, VITA (Fu et al., 2024b) employs TTS tools (RVC-Boss) to convert the output text into corresponding speech, while ModelVerse (Wang et al., 2024c) and UnifiedM-LLM (Li et al., 2024g) use the text to specify the generation model and utilize the corresponding descriptions to generate different modalities.

Modality-Token-based Works like MiniOmni-2 (Xie and Wu, 2024) and AnyGPT (Zhan et al., 2024) extend the corresponding LLM head with codebooks from different modality tokenizers to generate modality-specific discrete tokens. These tokens are then decoded using the corresponding de-tokenizers (Esser et al., 2021; Yu et al., 2024b; Zeghidour et al., 2022; Dhariwal et al., 2020) to produce various modalities.

Representation-based To alleviate the potential noise introduced by discrete tokens, works like X-VILA (Ye et al., 2024a) and NextGPT (Wu et al., 2024b) incorporate modality-specific signal tokens into the vocabulary. They then use transformers or MLPs to map the signal token representations into the ones that are understandable to the multimodal decoders, typically off-the-shelf latent-conditioned diffusion models (Rombach et al., 2022; Tang et al., 2023; Xue et al., 2024a; Blattmann et al., 2023), enabling effective generation capabilities.

3 Omni-MLLM Training

To achieve alignment across different vector spaces and improve instruction-following ability under arbitrary modality settings, Omni-MLLMs extend the standard two-stage training pipeline of Specific-MLLMs: *multi-modalities alignment pre-training* and *multi-modalities instruction fine-tuning*.

3.1 Multi-modalities Alignment Pre-training

Multi-modalities alignment pre-training involves *input alignment* training between the feature spaces of different modalities and the embedding space of LLMs on the encoding side, as well as *output alignment* training between the embedding space and the input spaces of various modality decoders on the decoding side. Input alignment and output alignment can be carried out separately (Wu et al., 2024b) or simultaneously (Ye et al., 2024a).

3.1.1 Input Alignment

Input alignment mainly uses X-Text paired datasets of different modalities and minimizes the text generation loss of the corresponding description text to optimize. In this phase, continuous encoding Omni-MLLMs normally update parameters of projectors, while discrete encoding Omni-MLLMs adjust the parameters of the embedding layer.

In terms of *training order* of different modalities alignment, multi-branch Omni-MLLM performs separate alignment training for each modality-specific projector, directly aligning each non-linguistic modality with text and using text as a bridge to align different non-linguistic modalities (Zhao et al., 2023b; Panagopoulou et al., 2024). The uni-branch Omni-MLLM, on the other hand, uses the unified projector for different modalities, which may lead to interference in the alignment performance between different modalities. Thus, Han et al. (2024a) employ a progressive alignment strategy to align multiple modalities in a specific order. In contrast, discrete encoding Omni-MLLMs, like AnyGPT (Zhan et al., 2024) and M3GPT (Luo et al., 2024), mix the alignment data from different modalities and perform alignment simultaneously.

Besides, in addition to directly leveraging X-Text paired datasets from different modalities for direct alignment, PandaGPT (Su et al., 2023), ImageBind-LLM (Han et al., 2023), and VideoL-LaMA (Zhang et al., 2023b) utilize the pre-aligned modality feature space $\mathbb{V}_{uni}$ to achieve indirect alignment between other non-linguistic modalities and text by training solely on Image-Text data.

3.1.2 Output Alignment

The training of output alignment typically utilizes the same X-text paired dataset as input alignment and adheres to the identical training sequence. Meanwhile, the training objectives for output alignment vary depending on the multi-modality generation methods in Section 2.4. Token-based gener-

ative Omni-MLLMs optimize the extended LLM head by minimizing the text generation loss associated with modality-specific discrete tokens (Lu et al., 2024a; Wei et al., 2024). Representation-based generative Omni-MLLMs generally optimize their output projectors by minimizing the composite loss comprising three components (Xie and Wu, 2024; Yang et al., 2023b): 1) the text generation loss of signal tokens; 2) the L2 distance between the output representation and the condition vector of the corresponding decoder, i.e. MSE loss; and 3) the conditional latent denoising loss (Rom-bach et al., 2022). For text-based generative Omni-MLLMs, as there is no additional output structure, the output alignment training is generally not required (Wang et al., 2024c; Li et al., 2024g).

3.2 Multi-modalities Instruction Fine-tuning

The instruction fine-tuning phase aims to enhance generalization capability under arbitrary modalities of Omni-MLLMs (Panagopoulou et al., 2024; Wu et al., 2024b; Ye et al., 2024a). Instruction fine-tuning primarily utilizes instruction-following datasets and computes the text generation loss for the corresponding responses to optimize. For models with generation capabilities, the loss mentioned in section 3.1.2 may also be incorporated. During this phase, Omni-MLLMs further perform full-scale tuning of the LLM parameters (Han et al., 2024a; Cheng et al., 2024b) or use PEFT techniques (Han et al., 2024b), such as LoRA (Hu et al., 2022), for partial tuning (Wang et al., 2024e).

Compared to Specific-MLLMs, Omni-MLLMs not only leverage multiple uni-modal instruction data of different modalities for training but also use cross-modal instruction data to enhance their cross-modal ability (Ye et al., 2024a; Zhan et al., 2024; Li et al., 2024c). In addition to directly mixing different instruction data for training (Panagopoulou et al., 2024; Ye et al., 2024a), some works like Uni-Moe (Li et al., 2024f) and Lyra (Zhong et al., 2024) adopt a multi-step fine-tuning approach, introducing different uni-modal and cross-modal instruction data in a specific order for training to gradually enhance their uni-modal and cross-modal ability.

4 Data Construction and Evaluation

This section summarizes the construction of modality alignment data and instruction data used in the Omni-MLLM training process (§4.1), as well as the evaluation across four different capabilities (§4.2).

4.1 Training Data

Alignment Data Omni-MLLMs leverage caption datasets from various modalities to construct X-Text paired data for alignment pre-training, such as the WebVid (Bain et al., 2021) for visual modality and the AudioCaps (Kim et al., 2019) for auditory modality. However, for data-scarce modalities like depth maps and thermal maps, large-scale text-paired data is lacking (Zhu et al., 2024a; Girdhar et al., 2023). To address this, synthetic methods that use DPT models (Ranftl et al., 2021; Bhat et al., 2023; Xu et al., 2023) or image translation models (Lee et al., 2023) to convert image-text pairs into other modality text pairs are widely employed (Han et al., 2024a; Chen et al., 2024c; Zhu et al., 2024a). Moreover, interleaved datasets (Zhu et al., 2023a) are used for alignment pre-training in some works (Tang et al., 2024b) to enhance the contextual understanding capability.

Instruction Data Omni-MLLMs not only leverage uni-modal instruction datasets from Specific-MLLMs, but also construct cross-modal instruction data through diverse methods as follows.

(1) **Template-based Construction:** Most works (Sun et al., 2023a; Zhao et al., 2023b; Zhang et al., 2024b) utilize cross-modal downstream datasets (Sanabria et al., 2018; Chen et al., 2020b) combined with predefined templates to construct cross-modal instructions; (2) **GPT Generation:** Following the paradigm of LLaVA (Liu et al., 2023c), some Omni-MLLMs (Lyu et al., 2023; Zhao et al., 2023b) leverage the labels from the annotated dataset (Lin et al., 2014; Bain et al., 2021) or use pre-trained models like SAM (Chen et al., 2023c) and GRIT (Wu et al., 2024a) to extract meta-information (e.g., captions and object categories) of different modalities. Then they employ powerful LLMs (OpenAI, 2023b,a) to generate cross-modal instructions based on the obtained meta-information; (3) **T2X Generation:** Li et al. (2024f) use TTS tools to convert the Image-Text2Text uni-modal instructions from LLaVA-v1.5 (Liu et al., 2024a) into Image-Speech-Text2Text cross-modal instructions. AnyGPT (Zhan et al., 2024) and NextGPT (Wu et al., 2024b) leverage Text2X models such as DALL-E-3 (Shi et al., 2020) and MusicGen (Copet et al., 2023) to convert the GPT-generated pure text instructions into Xs2Xs cross-modal instructions. Details about training data are shown in Appendix C.1

4.2 Benchmark

We provide a brief overview of the benchmarks used to evaluate Omni-MLLMs. The statistics of the benchmarks are shown in Appendix C.2.

Uni-modal Understanding Uni-modal understanding assesses the ability of Omni-MLLMs to comprehend and reason on different non-linguistic modalities, including downstream X-Text2Text datasets such as X-Caption (Plummer et al., 2015; Xu et al., 2016), X-QA (Goyal et al., 2017; Xu et al., 2017), and X-Classification (Deitke et al., 2023), as well as comprehensive multi-task benchmarks (Liu et al., 2024d; Fu et al., 2024a, 2023).

Uni-modal Generation Uni-modal generation aims to evaluate the ability of Omni-MLLMs to generate a single non-linguistic modality, including the Text2X generation task (Kim et al., 2019; Ruiz et al., 2023) and the Text-X2Text editing task (Veaux et al., 2017; Perazzi et al., 2016).

Cross-modal Understanding Cross-modal understanding evaluates the ability of Omni-MLLMs to jointly comprehend and reason across multiple non-linguistic modalities like Image-Speech-Text2Text (Li et al., 2024f; OpenGV, 2024), Video-Audio-Text2Text (Li et al., 2022a,b), and Image-3D-Text2Text (Panagopoulou et al., 2024).

Cross-modal Generation Cross-modal generation further evaluates the ability of Omni-MLLMs to generate non-linguistic modalities in conjunction with other non-linguistic modality inputs. For example, the Xs-Text2X benchmark proposed by X-VILA (Ye et al., 2024a) includes tasks such as Image-Text2Audio and Image-Audio-Text2Video.

5 Challenges and Future Directions

Despite Omni-MLLMs having showcased remarkable performance on numerous tasks, there are still some challenges that necessitate further research.

5.1 Expansion of modalities

Most Omni-MLLMs can only process 2-3 types of non-linguistic modalities, and they still face several challenges when expanding more modalities.

Training efficiency The common method that introduces new modalities through additional alignment pre-training and instruction fine-tuning can lead to significant training cost. Leveraging prior knowledge from Specific-MLLMs (Panagopoulou

et al., 2024; Chen et al., 2024a) or using pre-aligned encoders for indirect alignment (Han et al., 2023; Su et al., 2023) can help reduce training overhead but may impact cross-modal performance.

Catastrophic forgetting Expanding new modalities may adjust the shared parameters, potentially causing catastrophic forgetting of previously trained modalities knowledge (Yu et al., 2024a). This issue can be partially mitigated by mixing trained modality data (Han et al., 2024a; Li et al., 2024f) or fine-tuning only the modality-specific parameters (Yu et al., 2024a,c), but both approaches make the training process more complex.

Low-resource modalities Although the data synthesis method in Section 4.1 can help alleviate the lack of text-paired data and instruction data for low-resource modalities (Han et al., 2024a; ?), the absence of real modality may lead to biases in understanding of that modality.

5.2 Cross-modal capabilities

The Omni-MLLMs have achieved promising performance in cross-modal understanding and generation tasks, but there are still some challenges.

Long Context When the input contains multiple sequence modalities (video, speech...), the length of the multi-modalities token sequence may exceed the context window of LLMs and lead to memory overflow. While methods such as token compressing (Yu et al., 2024c; Li et al., 2024c) or token sampling (Zhan et al., 2024; Zhong et al., 2024) can reduce the number of input tokens, they also result in a decline in cross-modal performance.

Modality Bias Due to the imbalance in training data volume and the performance disparity among different modality encoders, Omni-MLLMs may tend to pay attention to the dominant modality while neglecting information from other modalities during cross-modal inference. Balancing the data volume across modalities or enhancing the corresponding modality-specific modules could potentially help mitigate this issue (Leng et al., 2024).

Temporal Alignment When dealing with different modalities that have temporal dependencies, retaining their temporal alignment information is crucial for subsequent cross-modal understanding. Some attempts have been made to preserve the temporal alignment information between audio and video, such as interleaved modality-specific tokens

of video and audio (Tang et al., 2024a) and inserting the time-related special tokens into the multi-modalities tokens (Goel et al., 2024).

Data and Benchmark Although Omni-MLLMs employ various methods in Section 4.1 to generate cross-modal instruction data, there is still significant room for improvement and expansion, including enhancing the diversity of instructions, incorporating longer contextual dialogues, and exploring more diverse modality interaction paradigms. Similarly, cross-modal benchmarks such as OmniBench (Li et al., 2024e) and OmniR (Chen et al., 2024d) still fall short in terms of task richness and instruction diversity when compared to uni-modal benchmarks like MMMU (Yue et al., 2024) and MME (Fu et al., 2023). And the variety of modalities they cover is also relatively limited.

5.3 Application scenarios

The emergence of Omni-MLLM brings new opportunities and possibilities for various applications. (1) **Real-time Multi-modalities Interaction:** Fu et al. (2025) and Xie and Wu (2024) achieve robust capabilities in both vision and speech understanding, enabling efficient speech-to-speech interactions with vision in real-time. (2) **Comprehensive Planning:** Wang et al. (2023b) and Szot et al. (2024a) leverage the complementarity across multiple modalities to achieve better path planning and action planning capabilities than planning with vision information only. (3) **World Simulator:** Ge et al. (2024b) not only understands and generates different modalities but also predicts state transitions for any combination of modalities.

6 Conclusion

In this paper, we provide a comprehensive survey report on Omni-MLLM, offering a comprehensive review of the field. Specifically, we break down Omni-MLLM into four key components and categorize them based on modal encoding and alignment methods. Subsequently, we provide a detailed summary of the training process of Omni-MLLM and the related resources used. We also summarize the current challenges and the future development directions. This paper is the first systematic survey dedicated to Omni-MLLMs. We hope this survey will facilitate further research in this area.

Limitations

This study provides the first comprehensive survey of Omni-MLLMs. Related work, architecture statistics, more details of training and evaluation, as well as other training recipes, can be found in Appendix A,B,C.

We have made our best effort, but there may still be some limitations. On one hand, due to page limitations, we can only provide a concise overview of the core contributions of mainstream Omni-MLLMs, rather than exhaustive technical details. On the other hand, our review primarily covers research from *ACL, NeurIPS, ICLR, ICML, COLING, CVPR, IJCAI, ECCV, and arXiv, and there is a chance that we may have missed some important work published in other venues. We will stay updated with ongoing discussions in the research community and plan to revise our work in the future to include overlooked contributions.

References

Panos Achlioptas, Ahmed Abdelreheem, Fei Xia, Mohamed Elhoseiny, and Leonidas J. Guibas. 2020. [Referit3d: Neural listeners for fine-grained 3d object identification in real-world scenes](#). In *Computer Vision - ECCV 2020 - 16th European Conference, Glasgow, UK, August 23-28, 2020, Proceedings, Part I*, volume 12346 of *Lecture Notes in Computer Science*, pages 422–440. Springer.

Andrea Agostinelli, Timo I. Denk, Zalán Borsos, Jesse H. Engel, Mauro Verzetti, Antoine Caillon, Qingqing Huang, Aren Jansen, Adam Roberts, Marco Tagliasacchi, Matthew Sharifi, Neil Zeghidour, and Christian Havnø Frank. 2023. [Musiclm: Generating music from text](#). *CoRR*, abs/2301.11325.

Harsh Agrawal, Peter Anderson, Karan Desai, Yufei Wang, Xinlei Chen, Rishabh Jain, Mark Johnson, Dhruv Batra, Devi Parikh, and Stefan Lee. 2019. [nocaps: novel object captioning at scale](#). In *2019 IEEE/CVF International Conference on Computer Vision, ICCV 2019, Seoul, Korea (South), October 27 - November 2, 2019*, pages 8947–8956. IEEE.

Huda AlAmri, Vincent Cartillier, Raphael Gontijo Lopes, Abhishek Das, Jue Wang, Irfan Essa, Dhruv Batra, Devi Parikh, Anoop Cherian, Tim K. Marks, and Chiori Hori. 2018. [Audio visual scene-aware dialog \(AVSD\) challenge at DSTC7](#). *CoRR*, abs/1806.00525.

Emily J Allen, Ghislain St-Yves, Yihan Wu, Jesse L Breedlove, Jacob S Prince, Logan T Dowdle, Matthias Nau, Brad Caron, Franco Pestilli, Ian Charest, et al. 2022. A massive 7t fmri dataset to bridge cognitive neuroscience and artificial intelligence. *Nature neuroscience*, 25(1):116–126.

Anonymous. 2024. [Aligned better, listen better for audio-visual large language models](#). In *Submitted to The Thirteenth International Conference on Learning Representations*. Under review.

Rosana Ardila, Megan Branson, Kelly Davis, Michael Kohler, Josh Meyer, Michael Henretty, Reuben Morais, Lindsay Saunders, Francis M. Tyers, and Gregor Weber. 2020. [Common voice: A massively-multilingual speech corpus](#). In *Proceedings of The 12th Language Resources and Evaluation Conference, LREC 2020, Marseille, France, May 11-16, 2020*, pages 4218–4222. European Language Resources Association.

Anurag Arnab, Mostafa Dehghani, Georg Heigold, Chen Sun, Mario Lucic, and Cordelia Schmid. 2021. [Vivit: A video vision transformer](#). In *2021 IEEE/CVF International Conference on Computer Vision, ICCV 2021, Montreal, QC, Canada, October 10-17, 2021*, pages 6816–6826. IEEE.

Daichi Azuma, Taiki Miyanishi, Shuhei Kurita, and Motoaki Kawanabe. 2022. [Scanqa: 3d question answering for spatial scene understanding](#). In *IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2022, New Orleans, LA, USA, June 18-24, 2022*, pages 19107–19117. IEEE.

Fan Bai, Yuxin Du, Tiejun Huang, Max Qinghu Meng, and Bo Zhao. 2024a. [M3D: advancing 3d medical image analysis with multi-modal large language models](#). *CoRR*, abs/2404.00578.

Jinze Bai, Shuai Bai, Yunfei Chu, Zeyu Cui, Kai Dang, Xiaodong Deng, Yang Fan, Wenbin Ge, Yu Han, Fei Huang, Binyuan Hui, Luo Ji, Mei Li, Junyang Lin, Runji Lin, Dayiheng Liu, Gao Liu, Chengqiang Lu, Keming Lu, Jianxin Ma, Rui Men, Xingzhang Ren, Xuancheng Ren, Chuanqi Tan, Sinan Tan, Jianhong Tu, Peng Wang, Shijie Wang, Wei Wang, Shengguang Wu, Benfeng Xu, Jin Xu, An Yang, Hao Yang, Jian Yang, Shusheng Yang, Yang Yao, Bowen Yu, Hongyi Yuan, Zheng Yuan, Jianwei Zhang, Xingxuan Zhang, Yichang Zhang, Zhenru Zhang, Chang Zhou, Jingren Zhou, Xiaohuan Zhou, and Tianhang Zhu. 2023. [Qwen technical report](#). *CoRR*, abs/2309.16609.

Zeichen Bai, Pichao Wang, Tianjun Xiao, Tong He, Zongbo Han, Zheng Zhang, and Mike Zheng Shou. 2024b. [Hallucination of multimodal large language models: A survey](#). *CoRR*, abs/2404.18930.

Max Bain, Arsha Nagrani, Gül Varol, and Andrew Senior. 2021. [Frozen in time: A joint video and image encoder for end-to-end retrieval](#). In *2021 IEEE/CVF International Conference on Computer Vision, ICCV 2021, Montreal, QC, Canada, October 10-17, 2021*, pages 1708–1718. IEEE.

Gedas Bertasius, Heng Wang, and Lorenzo Torresani. 2021. [Is space-time attention all you need for video understanding?](#) In *Proceedings of the 38th International Conference on Machine Learning, ICML 2021, 18-24 July 2021, Virtual Event*, volume 139 of

762	<i>Proceedings of Machine Learning Research</i> , pages	Linyang Li, Shuaibin Li, Wei Li, Yining Li, Hong-	819
763	813–824. PMLR.	wei Liu, Jiangning Liu, Jiawei Hong, Kaiwen Liu,	820
764	Shariq Farooq Bhat, Reiner Birkel, Diana Wofk, Peter	Kuikun Liu, Xiaoran Liu, Chengqi Lv, Haijun Lv,	821
765	Wonka, and Matthias Müller. 2023. Zoedepth: Zero-	Kai Lv, Li Ma, Runyuan Ma, Zerun Ma, Wenchang	822
766	shot transfer by combining relative and metric depth.	Ning, Linke Ouyang, Jiantao Qiu, Yuan Qu, Fukai	823
767	<i>CoRR</i> , abs/2302.12288.	Shang, Yunfan Shao, Demin Song, Zifan Song, Zhi-	824
768	Ali Furkan Biten, Rubèn Tito, Andrés Mafla, Lluís	hao Sui, Peng Sun, Yu Sun, Huanze Tang, Bin Wang,	825
769	Gómez, Marçal Rusiñol, Minesh Mathew, C. V. Jawa-	Guoteng Wang, Jiaqi Wang, Jiayu Wang, Rui Wang,	826
770	har, Ernest Valveny, and Dimosthenis Karatzas. 2019.	Yudong Wang, Ziyi Wang, Xingjian Wei, Qizhen	827
771	ICDAR 2019 competition on scene text visual ques-	Weng, Fan Wu, Yingdong Xiong, Xiaomeng Zhao,	828
772	tion answering. In <i>2019 International Conference on</i>	and et al. 2024. Internlm2 technical report. <i>CoRR</i> ,	829
773	<i>Document Analysis and Recognition, ICDAR 2019,</i>	abs/2403.17297.	830
774	<i>Sydney, Australia, September 20-25, 2019</i> , pages	Umberto Cappellazzo, Minsu Kim, Honglie Chen,	831
775	1563–1570. IEEE.	Pingchuan Ma, Stavros Petridis, Daniele Falavigna,	832
776	Andreas Blattmann, Tim Dockhorn, Sumith Ku-	Alessio Brutti, and Maja Pantic. 2024. Large lan-	833
777	lal, Daniel Mendelevitch, Maciej Kilian, Dominik	guage models are strong audio-visual speech recog-	834
778	Lorenz, Yam Levi, Zion English, Vikram Voleti,	nition learners. <i>CoRR</i> , abs/2409.12319.	835
779	Adam Letts, Varun Jampani, and Robin Rom-	Cerspense. 2023. Zeroscope: Diffusion-based text-to-	836
780	bach. 2023. Stable video diffusion: Scaling latent	video synthesis.	837
781	video diffusion models to large datasets. <i>CoRR</i> ,	Junbum Cha, Wooyoung Kang, Jonghwan Mun, and	838
782	abs/2311.15127.	Byungseok Roh. 2024. Honeybee: Locality-	839
783	Tim Brooks, Aleksander Holynski, and Alexei A. Efros.	enhanced projector for multimodal LLM. In	840
784	2023. Instructpix2pix: Learning to follow image	<i>IEEE/CVF Conference on Computer Vision and Pat-</i>	841
785	editing instructions. In <i>IEEE/CVF Conference on</i>	<i>tern Recognition, CVPR 2024, Seattle, WA, USA,</i>	842
786	<i>Computer Vision and Pattern Recognition, CVPR</i>	<i>June 16-22, 2024</i> , pages 13817–13827. IEEE.	843
787	2023, <i>Vancouver, BC, Canada, June 17-24, 2023</i> ,	Soravit Changpinyo, Piyush Sharma, Nan Ding, and	844
788	pages 18392–18402. IEEE.	Radu Soricut. 2021. Conceptual 12m: Pushing web-	845
789	Hui Bu, Jiayu Du, Xingyu Na, Bengu Wu, and Hao	scale image-text pre-training to recognize long-tail	846
790	Zheng. 2017. AISHELL-1: an open-source man-	dard visual concepts. In <i>IEEE Conference on Computer</i>	847
791	darin speech corpus and a speech recognition base-	<i>Vision and Pattern Recognition, CVPR 2021, virtual,</i>	848
792	line. In <i>20th Conference of the Oriental Chapter of</i>	<i>June 19-25, 2021</i> , pages 3558–3568. Computer Vi-	849
793	<i>the International Coordinating Committee on Speech</i>	sion Foundation / IEEE.	850
794	<i>Databases and Speech I/O Systems and Assessment,</i>	Chi Chen, Yiyang Du, Zheng Fang, Ziyue Wang, Fuwen	851
795	<i>O-COCOSDA 2017, Seoul, South Korea, November</i>	Luo, Peng Li, Ming Yan, Ji Zhang, Fei Huang,	852
796	<i>1-3, 2017</i> , pages 1–5. IEEE.	Maosong Sun, and Yang Liu. 2024a. Model com-	853
797	Sébastien Bubeck, Varun Chandrasekaran, Ronen Eldan,	position for multimodal large language models. In	854
798	Johannes Gehrmann, Eric Horvitz, Ece Kamar, Peter	<i>Proceedings of the 62nd Annual Meeting of the As-</i>	855
799	Lee, Yin Tat Lee, Yuanzhi Li, Scott M. Lundberg,	<i>sociation for Computational Linguistics (Volume 1:</i>	856
800	Harsha Nori, Hamid Palangi, Marco Túlio Ribeiro,	<i>Long Papers)</i> , <i>ACL 2024, Bangkok, Thailand, August</i>	857
801	and Yi Zhang. 2023. Sparks of artificial general	<i>11-16, 2024</i> , pages 11246–11262. Association for	858
802	intelligence: Early experiments with GPT-4. <i>CoRR</i> ,	Computational Linguistics.	859
803	abs/2303.12712.	Dave Zhenyu Chen, Angel X. Chang, and Matthias	860
804	Davide Caffagni, Federico Cocchi, Luca Barsellotti,	Nießner. 2020a. Scanrefer: 3d object localization in	861
805	Nicholas Moratelli, Sara Sarto, Lorenzo Baraldi, Mar-	RGB-D scans using natural language. In <i>Computer</i>	862
806	cella Cornia, and Rita Cucchiara. 2024. The revolu-	<i>Vision - ECCV 2020 - 16th European Conference,</i>	863
807	tion of multimodal large language models: A survey.	<i>Glasgow, UK, August 23-28, 2020, Proceedings, Part</i>	864
808	In <i>Findings of the Association for Computational Lin-</i>	<i>XX</i> , volume 12365 of <i>Lecture Notes in Computer</i>	865
809	<i>guistics, ACL 2024, Bangkok, Thailand and virtual</i>	<i>Science</i> , pages 202–221. Springer.	866
810	<i>meeting, August 11-16, 2024</i> , pages 13590–13618.	Feilong Chen, Minglun Han, Haozhi Zhao, Qingyang	867
811	Association for Computational Linguistics.	Zhang, Jing Shi, Shuang Xu, and Bo Xu. 2023a. X-	868
812	Zheng Cai, Maosong Cao, Haojiong Chen, Kai Chen,	LLM: bootstrapping advanced large language models	869
813	Keyu Chen, Xin Chen, Xun Chen, Zehui Chen, Zhi	by treating multi-modalities as foreign languages.	870
814	Chen, Pei Chu, Xiaoyi Dong, Haodong Duan, Qi Fan,	<i>CoRR</i> , abs/2305.04160.	871
815	Zhaoye Fei, Yang Gao, Jiaye Ge, Chenya Gu, Yuzhe	Guoguo Chen, Shuzhou Chai, Guan-Bo Wang, Jiayu	872
816	Gu, Tao Gui, Aijia Guo, Qipeng Guo, Conghui He,	Du, Wei-Qiang Zhang, Chao Weng, Dan Su, Daniel	873
817	Yingfan Hu, Ting Huang, Tao Jiang, Penglong Jiao,	Povey, Jan Trmal, Junbo Zhang, Mingjie Jin, San-	874
818	Zhenjiang Jin, Zhikai Lei, Jiaying Li, Jingwen Li,	jeev Khudanpur, Shinji Watanabe, Shuaijiang Zhao,	875

876	Wei Zou, Xiangang Li, Xuchen Yao, Yongqing Wang,	202 of <i>Proceedings of Machine Learning Research</i> ,	933
877	Zhao You, and Zhiyong Yan. 2021. Gigaspeech: An	pages 5178–5193. PMLR.	934
878	evolving, multi-domain ASR corpus with 10, 000		
879	hours of transcribed audio . In <i>22nd Annual Con-</i>	Sihan Chen, Xingjian He, Longteng Guo, Xinxin Zhu,	935
880	<i>ference of the International Speech Communication</i>	Weining Wang, Jinhui Tang, and Jing Liu. 2023e.	936
881	<i>Association, Interspeech 2021, Brno, Czechia, August</i>	VALOR: vision-audio-language omni-perception pre-	937
882	<i>30 - September 3, 2021</i> , pages 3670–3674. ISCA.	training model and dataset . <i>CoRR</i> , abs/2304.08345.	938
883	Haoxin Chen, Menghan Xia, Yingqing He, Yong Zhang,	Sihan Chen, Handong Li, Qunbo Wang, Zijia Zhao,	939
884	Xiaodong Cun, Shaoshu Yang, Jinbo Xing, Yaofang	Mingzhen Sun, Xinxin Zhu, and Jing Liu. 2023f.	940
885	Liu, Qifeng Chen, Xintao Wang, Chao Weng, and	VAST: A vision-audio-subtitle-text omni-modality	941
886	Ying Shan. 2023b. Videocrafter1: Open diffusion	foundation model and dataset . In <i>Advances in Neural</i>	942
887	models for high-quality video generation . <i>CoRR</i> ,	<i>Information Processing Systems 36: Annual Confer-</i>	943
888	abs/2310.19512.	<i>ence on Neural Information Processing Systems 2023,</i>	944
889	Hong Chen, Xin Wang, Yuwei Zhou, Bin Huang,	<i>NeurIPS 2023, New Orleans, LA, USA, December 10</i>	945
890	Yipeng Zhang, Wei Feng, Houlun Chen, Zeyang	<i>- 16, 2023</i> .	946
891	Zhang, Siao Tang, and Wenwu Zhu. 2024b. Multi-	Tsai-Shien Chen, Aliaksandr Siarohin, Willi Menapace,	947
892	modal generative AI: multi-modal llm, diffusion and	Ekaterina Deyneka, Hsiang-wei Chao, Byung Eun	948
893	beyond . <i>CoRR</i> , abs/2409.14993.	Jeon, Yuwei Fang, Hsin-Ying Lee, Jian Ren, Ming-	949
894	Honglie Chen, Weidi Xie, Andrea Vedaldi, and An-	Hsuan Yang, and Sergey Tulyakov. 2024f. Panda-	950
895	drew Zisserman. 2020b. Vggsound: A large-scale	70m: Captioning 70m videos with multiple cross-	951
896	audio-visual dataset . In <i>2020 IEEE International</i>	modality teachers . In <i>IEEE/CVF Conference on Com-</i>	952
897	<i>Conference on Acoustics, Speech and Signal Process-</i>	<i>puter Vision and Pattern Recognition, CVPR 2024,</i>	953
898	<i>ing, ICASSP 2020, Barcelona, Spain, May 4-8, 2020,</i>	<i>Seattle, WA, USA, June 16-22, 2024</i> , pages 13320–	954
899	pages 721–725. IEEE.	13331. IEEE.	955
900	Jiaqi Chen, Zeyu Yang, and Li Zhang. 2023c. Se-	Zeren Chen, Ziqin Wang, Zhen Wang, Huayang Liu,	956
901	mantic segment anything. https://github.com/	Zhenfei Yin, Si Liu, Lu Sheng, Wanli Ouyang, and	957
902	fudan-zvg/Semantic-Segment-Anything .	Jing Shao. 2024g. Octavius: Mitigating task inter-	958
903	Kai Chen, Yunhao Gou, Runhui Huang, Zhili Liu, Daxin	ference in mllms via lora-moe . In <i>The Twelfth In-</i>	959
904	Tan, Jing Xu, Chunwei Wang, Yi Zhu, Yihan Zeng,	<i>ternational Conference on Learning Representations,</i>	960
905	Kuo Yang, Dingdong Wang, Kun Xiang, Haoyuan Li,	<i>ICLR 2024, Vienna, Austria, May 7-11, 2024</i> . Open-	961
906	Haoli Bai, Jianhua Han, Xiaohui Li, WeiKe Jin, Nian	Review.net.	962
907	Xie, Yu Zhang, James T. Kwok, Hengshuang Zhao,	Zhe Chen, Jiannan Wu, Wenhai Wang, Weijie Su, Guo	963
908	Xiaodan Liang, Dit-Yan Yeung, Xiao Chen, Zhenguo	Chen, Sen Xing, Muyan Zhong, Qinglong Zhang,	964
909	Li, Wei Zhang, Qun Liu, Jun Yao, Lanqing Hong,	Xizhou Zhu, Lewei Lu, Bin Li, Ping Luo, Tong Lu,	965
910	Lu Hou, and Hang Xu. 2024c. EMOVA: empowering	Yu Qiao, and Jifeng Dai. 2023g. Internvl: Scaling	966
911	language models to see, hear and speak with vivid	up vision foundation models and aligning for generic	967
912	emotions . <i>CoRR</i> , abs/2409.18042.	visual-linguistic tasks . <i>CoRR</i> , abs/2312.14238.	968
913	Lichang Chen, Hexiang Hu, Mingda Zhang, Yiwen	Zebang Cheng, Zhi-Qi Cheng, Jun-Yan He, Jingdong	969
914	Chen, Zifeng Wang, Yandong Li, Pranav Shyam,	Sun, Kai Wang, Yuxiang Lin, Zheng Lian, Xiaojiang	970
915	Tianyi Zhou, Heng Huang, Ming-Hsuan Yang, and	Peng, and Alexander Hauptmann. 2024a. Emotion-	971
916	Boqing Gong. 2024d. Omnixr: Evaluating omni-	llama: Multimodal emotion recognition and reason-	972
917	modality language models on reasoning across	ing with instruction tuning . <i>CoRR</i> , abs/2406.11161.	973
918	modalities . <i>CoRR</i> , abs/2410.12219.	Zesen Cheng, Sicong Leng, Hang Zhang, Yifei Xin, Xin	974
919	Lin Chen, Jinsong Li, Xiaoyi Dong, Pan Zhang, Cong-	Li, Guanzheng Chen, Yongxin Zhu, Wenqi Zhang,	975
920	hui He, Jiaqi Wang, Feng Zhao, and Dahua Lin.	Ziyang Luo, Deli Zhao, and Lidong Bing. 2024b.	976
921	2024e. Sharegpt4v: Improving large multi-modal	Videollama 2: Advancing spatial-temporal model-	977
922	models with better captions . In <i>Computer Vision</i>	ing and audio understanding in video-llms . <i>CoRR</i> ,	978
923	<i>- ECCV 2024 - 18th European Conference, Milan,</i>	abs/2406.07476.	979
924	<i>Italy, September 29-October 4, 2024, Proceedings,</i>	Chee Kheng Chng and Chee Seng Chan. 2017. Total-	980
925	<i>Part XVII</i> , volume 15075 of <i>Lecture Notes in Com-</i>	text: A comprehensive dataset for scene text detec-	981
926	<i>puter Science</i> , pages 370–387. Springer.	tion and recognition . In <i>14th IAPR International</i>	982
927	Sanyuan Chen, Yu Wu, Chengyi Wang, Shujie Liu,	<i>Conference on Document Analysis and Recognition,</i>	983
928	Daniel Tompkins, Zhuo Chen, Wanxiang Che, Xi-	<i>ICDAR 2017, Kyoto, Japan, November 9-15, 2017,</i>	984
929	angzhan Yu, and Furu Wei. 2023d. Beats: Audio	pages 935–942. IEEE.	985
930	pre-training with acoustic tokenizers . In <i>Interna-</i>	Sanjoy Chowdhury, Sayan Nag, Subhrajyoti Dasgupta,	986
931	<i>tional Conference on Machine Learning, ICML 2023,</i>	Jun Chen, Mohamed Elhoseiny, Ruohan Gao, and Di-	987
932	<i>23-29 July 2023, Honolulu, Hawaii, USA</i> , volume	nesh Manocha. 2024. MEERKAT: audio-visual large	988
		language model for grounding in space and time . In	989

990	<i>Computer Vision - ECCV 2024 - 18th European Conference, Milan, Italy, September 29-October 4, 2024, Proceedings, Part LXIV</i> , volume 15122 of <i>Lecture Notes in Computer Science</i> , pages 52–70. Springer.	
991		
992		
993		
994	Yunfei Chu, Jin Xu, Xiaohuan Zhou, Qian Yang, Shiliang Zhang, Zhijie Yan, Chang Zhou, and Jingren Zhou. 2023. Qwen-audio: Advancing universal audio understanding via unified large-scale audio-language models . <i>CoRR</i> , abs/2311.07919.	
995		
996		
997		
998		
999	Enis Berk Çoban, Michael I. Mandel, and Johanna Devaney. 2024. What do mllms hear? examining reasoning with text and sound components in multimodal large language models . <i>CoRR</i> , abs/2406.04615.	
1000		
1001		
1002		
1003	Jade Copet, Felix Kreuk, Itai Gat, Tal Remez, David Kant, Gabriel Synnaeve, Yossi Adi, and Alexandre Défossez. 2023. Simple and controllable music generation . In <i>Advances in Neural Information Processing Systems 36: Annual Conference on Neural Information Processing Systems 2023, NeurIPS 2023, New Orleans, LA, USA, December 10 - 16, 2023</i> .	
1004		
1005		
1006		
1007		
1008		
1009		
1010	Can Cui, Yunsheng Ma, Xu Cao, Wenqian Ye, Yang Zhou, Kaizhao Liang, Jintai Chen, Juanwu Lu, Zichong Yang, Kuei-Da Liao, Tianren Gao, Erlong Li, Kun Tang, Zhipeng Cao, Tong Zhou, Ao Liu, Xinrui Yan, Shuqi Mei, Jianguo Cao, Ziran Wang, and Chao Zheng. 2024. A survey on multimodal large language models for autonomous driving . In <i>IEEE/CVF Winter Conference on Applications of Computer Vision Workshops, WACVW 2024 - Workshops, Waikoloa, HI, USA, January 1-6, 2024</i> , pages 958–979. IEEE.	
1011		
1012		
1013		
1014		
1015		
1016		
1017		
1018		
1019		
1020	Abhishek Das, Satwik Kottur, Khushi Gupta, Avi Singh, Deshraj Yadav, José M. F. Moura, Devi Parikh, and Dhruv Batra. 2017. Visual dialog . In <i>2017 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2017, Honolulu, HI, USA, July 21-26, 2017</i> , pages 1080–1089. IEEE Computer Society.	
1021		
1022		
1023		
1024		
1025		
1026	Alexandre Défossez, Jade Copet, Gabriel Synnaeve, and Yossi Adi. 2023. High fidelity neural audio compression . <i>Trans. Mach. Learn. Res.</i> , 2023.	
1027		
1028		
1029	Matt Deitke, Dustin Schwenk, Jordi Salvador, Luca Weihs, Oscar Michel, Eli VanderBilt, Ludwig Schmidt, Kiana Ehsani, Aniruddha Kembhavi, and Ali Farhadi. 2023. Objaverse: A universe of annotated 3d objects . In <i>IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2023, Vancouver, BC, Canada, June 17-24, 2023</i> , pages 13142–13153. IEEE.	
1030		
1031		
1032		
1033		
1034		
1035		
1036		
1037	Karan Desai, Gaurav Kaul, Zubin Aysola, and Justin Johnson. 2021. Redcaps: Web-curated image-text data created by the people, for the people . In <i>Proceedings of the Neural Information Processing Systems Track on Datasets and Benchmarks 1, NeurIPS Datasets and Benchmarks 2021, December 2021, virtual</i> .	
1038		
1039		
1040		
1041		
1042		
1043		
1044	Prafulla Dhariwal, Heewoo Jun, Christine Payne, Jong Wook Kim, Alec Radford, and Ilya Sutskever. 2020. Jukebox: A generative model for music . <i>CoRR</i> , abs/2005.00341.	
1045		
1046		
1047		
	Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, Jakob Uszkoreit, and Neil Houlsby. 2021. An image is worth 16x16 words: Transformers for image recognition at scale . In <i>9th International Conference on Learning Representations, ICLR 2021, Virtual Event, Austria, May 3-7, 2021</i> . OpenReview.net.	1048
		1049
		1050
		1051
		1052
		1053
		1054
		1055
		1056
	Konstantinos Drossos, Samuel Lipping, and Tuomas Virtanen. 2020. Clotho: an audio captioning dataset . In <i>2020 IEEE International Conference on Acoustics, Speech and Signal Processing, ICASSP 2020, Barcelona, Spain, May 4-8, 2020</i> , pages 736–740. IEEE.	1057
		1058
		1059
		1060
		1061
		1062
	Jiayu Du, Xingyu Na, Xuechen Liu, and Hui Bu. 2018. AISHELL-2: transforming mandarin ASR research into industrial scale . <i>CoRR</i> , abs/1808.10583.	1063
		1064
		1065
	Benjamin Elizalde, Soham Deshmukh, Mahmoud Al Ismail, and Huaming Wang. 2023. CLAP learning audio concepts from natural language supervision . In <i>IEEE International Conference on Acoustics, Speech and Signal Processing ICASSP 2023, Rhodes Island, Greece, June 4-10, 2023</i> , pages 1–5. IEEE.	1066
		1067
		1068
		1069
		1070
		1071
	Patrick Esser, Robin Rombach, and Björn Ommer. 2021. Taming transformers for high-resolution image synthesis . In <i>IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2021, virtual, June 19-25, 2021</i> , pages 12873–12883. Computer Vision Foundation / IEEE.	1072
		1073
		1074
		1075
		1076
		1077
	Yihe Fan, Yuxin Cao, Ziyu Zhao, Ziyao Liu, and Shaofeng Li. 2024. Unbridled icarus: A survey of the potential perils of image inputs in multimodal large language model security . <i>CoRR</i> , abs/2404.05264.	1078
		1079
		1080
		1081
	Alex Fang, Albin Madappally Jose, Amit Jain, Ludwig Schmidt, Alexander T. Toshev, and Vaishaal Shankar. 2024a. Data filtering networks . In <i>The Twelfth International Conference on Learning Representations, ICLR 2024, Vienna, Austria, May 7-11, 2024</i> . OpenReview.net.	1082
		1083
		1084
		1085
		1086
		1087
	Qingkai Fang, Shoutao Guo, Yan Zhou, Zhengrui Ma, Shaolei Zhang, and Yang Feng. 2024b. Llama-omni: Seamless speech interaction with large language models . <i>CoRR</i> , abs/2409.06666.	1088
		1089
		1090
		1091
	Hao Fei, Han Zhang, Bin Wang, Lizi Liao, Qian Liu, and Erik Cambria. 2024. Empathyear: An open-source avatar multimodal empathetic chatbot . <i>CoRR</i> , abs/2406.15177.	1092
		1093
		1094
		1095
	Chaoyou Fu, Peixian Chen, Yunhang Shen, Yulei Qin, Mengdan Zhang, Xu Lin, Zhenyu Qiu, Wei Lin, Jinrui Yang, Xiawu Zheng, Ke Li, Xing Sun, and Rongrong Ji. 2023. MME: A comprehensive evaluation benchmark for multimodal large language models . <i>CoRR</i> , abs/2306.13394.	1096
		1097
		1098
		1099
		1100
		1101

- Chaoyou Fu, Yuhan Dai, Yondong Luo, Lei Li, Shuhuai Ren, Renrui Zhang, Zihan Wang, Chenyu Zhou, Yunhang Shen, Mengdan Zhang, Peixian Chen, Yanwei Li, Shaohui Lin, Sirui Zhao, Ke Li, Tong Xu, Xiwu Zheng, Enhong Chen, Rongrong Ji, and Xing Sun. 2024a. [Video-mme: The first-ever comprehensive evaluation benchmark of multi-modal llms in video analysis](#). *CoRR*, abs/2405.21075.
- Chaoyou Fu, Haojia Lin, Zuwei Long, Yunhang Shen, Meng Zhao, Yifan Zhang, Xiong Wang, Di Yin, Long Ma, Xiwu Zheng, Ran He, Rongrong Ji, Yunsheng Wu, Caifeng Shan, and Xing Sun. 2024b. [VITA: towards open-source interactive omni multimodal LLM](#). *CoRR*, abs/2408.05211.
- Chaoyou Fu, Haojia Lin, Xiong Wang, Yi-Fan Zhang, Yunhang Shen, Xiaoyu Liu, Yangze Li, Zuwei Long, Heting Gao, Ke Li, Long Ma, Xiwu Zheng, Rongrong Ji, Xing Sun, Caifeng Shan, and Ran He. 2025. [Vita-1.5: Towards gpt-4o level real-time vision and speech interaction](#). *Preprint*, arXiv:2501.01957.
- Letian Fu, Gaurav Datta, Huang Huang, William Chung-Ho Panitch, Jaimyn Drake, Joseph Ortiz, Mustafa Mukadam, Mike Lambeta, Roberto Calandra, and Ken Goldberg. 2024c. [A touch, vision, and language dataset for multimodal alignment](#). In *Forty-first International Conference on Machine Learning, ICML 2024, Vienna, Austria, July 21-27, 2024*. OpenReview.net.
- Yuying Ge, Yixiao Ge, Ziyun Zeng, Xintao Wang, and Ying Shan. 2023. [Planting a SEED of vision in large language model](#). *CoRR*, abs/2307.08041.
- Yuying Ge, Sijie Zhao, Ziyun Zeng, Yixiao Ge, Chen Li, Xintao Wang, and Ying Shan. 2024a. [Making llama SEE and draw with SEED tokenizer](#). In *The Twelfth International Conference on Learning Representations, ICLR 2024, Vienna, Austria, May 7-11, 2024*. OpenReview.net.
- Zhiqi Ge, Hongzhe Huang, Mingze Zhou, Juncheng Li, Guoming Wang, Siliang Tang, and Yueting Zhuang. 2024b. [Worldgpt: Empowering LLM as multimodal world model](#). In *Proceedings of the 32nd ACM International Conference on Multimedia, MM 2024, Melbourne, VIC, Australia, 28 October 2024 - 1 November 2024*, pages 7346–7355. ACM.
- Jort F. Gemmeke, Daniel P. W. Ellis, Dylan Freedman, Aren Jansen, Wade Lawrence, R. Channing Moore, Manoj Plakal, and Marvin Ritter. 2017. [Audio set: An ontology and human-labeled dataset for audio events](#). In *2017 IEEE International Conference on Acoustics, Speech and Signal Processing, ICASSP 2017, New Orleans, LA, USA, March 5-9, 2017*, pages 776–780. IEEE.
- Tiantian Geng, Teng Wang, Jinming Duan, Runmin Cong, and Feng Zheng. 2023. [Dense-localizing audio-visual events in untrimmed videos: A large-scale benchmark and baseline](#). In *IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2023, Vancouver, BC, Canada, June 17-24, 2023*, pages 22942–22951. IEEE.
- Tiantian Geng, Jinrui Zhang, Qingni Wang, Teng Wang, Jinming Duan, and Feng Zheng. 2024. [Longvale: Vision-audio-language-event benchmark towards time-aware omni-modal perception of long videos](#). *CoRR*, abs/2411.19772.
- Rohit Girdhar, Alaaeldin El-Nouby, Zhuang Liu, Manan Singh, Kalyan Vasudev Alwala, Armand Joulin, and Ishan Misra. 2023. [Imagebind one embedding space to bind them all](#). In *IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2023, Vancouver, BC, Canada, June 17-24, 2023*, pages 15180–15190. IEEE.
- Arushi Goel, Karan Sapra, Matthieu Le, Rafael Valle, Andrew Tao, and Bryan Catanzaro. 2024. [OM-CAT: omni context aware transformer](#). *CoRR*, abs/2410.12109.
- Yuan Gong, Yu-An Chung, and James R. Glass. 2021. [AST: audio spectrogram transformer](#). In *22nd Annual Conference of the International Speech Communication Association, Interspeech 2021, Brno, Czechia, August 30 - September 3, 2021*, pages 571–575. ISCA.
- Yuan Gong, Jin Yu, and James R. Glass. 2022. [Vocal-sound: A dataset for improving human vocal sounds recognition](#). In *IEEE International Conference on Acoustics, Speech and Signal Processing, ICASSP 2022, Virtual and Singapore, 23-27 May 2022*, pages 151–155. IEEE.
- Yash Goyal, Tejas Khot, Douglas Summers-Stay, Dhruv Batra, and Devi Parikh. 2017. [Making the V in VQA matter: Elevating the role of image understanding in visual question answering](#). In *2017 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2017, Honolulu, HI, USA, July 21-26, 2017*, pages 6325–6334. IEEE Computer Society.
- Kristen Grauman, Andrew Westbury, Eugene Byrne, Zachary Chavis, Antonino Furnari, Rohit Girdhar, Jackson Hamburger, Hao Jiang, Miao Liu, Xingyu Liu, Miguel Martin, Tushar Nagarajan, Ilija Radosavovic, Santhosh Kumar Ramakrishnan, Fiona Ryan, Jayant Sharma, Michael Wray, Mengmeng Xu, Eric Zhongcong Xu, Chen Zhao, Siddhant Bansal, Dhruv Batra, Vincent Cartillier, Sean Crane, Tien Do, Morrie Doulaty, Akshay Erapalli, Christoph Feichtenhofer, Adriano Fragomeni, Qichen Fu, Abraham Gebreselasie, Cristina González, James Hillis, Xuhua Huang, Yifei Huang, Wenqi Jia, Weslie Khoo, Jáchym Kolár, Satwik Kottur, Anurag Kumar, Federico Landini, Chao Li, Yanghao Li, Zhenqiang Li, Karttikeya Mangalam, Raghava Modhugu, Jonathan Munro, Tullie Murrell, Takumi Nishiyasu, Will Price, Paola Ruiz Puentes, Merey Ramazanova, Leda Sari, Kiran Somasundaram, Audrey Southerland, Yusuke Sugano, Ruijie Tao, Minh Vo, Yuchen Wang, Xindi Wu, Takuma Yagi, Ziwei Zhao, Yunyi Zhu, Pablo Arbeláez, David Crandall, Dima Damen, Giovanni Maria Farinella, Christian Fuegen, Bernard

1218	Ghanem, Vamsi Krishna Ithapu, C. V. Jawahar, Han-	1275
1219	byul Joo, Kris Kitani, Haizhou Li, Richard A. New-	1276
1220	combe, Aude Oliva, Hyun Soo Park, James M. Rehg,	1277
1221	Yoichi Sato, Jianbo Shi, Mike Zheng Shou, Anto-	1278
1222	nio Torralba, Lorenzo Torresani, Mingfei Yan, and	
1223	Jitendra Malik. 2022. Ego4d: Around the world in	1279
1224	3, 000 hours of egocentric video . In <i>IEEE/CVF Con-</i>	1280
1225	<i>ference on Computer Vision and Pattern Recognition,</i>	1281
1226	<i>CVPR 2022, New Orleans, LA, USA, June 18-24,</i>	1282
1227	2022, pages 18973–18990. IEEE.	1283
		1284
1228	Jiayi Gu, Xiaojun Meng, Guansong Lu, Lu Hou, Niu	
1229	Minzhe, Xiaodan Liang, Lewei Yao, Runhui Huang,	1285
1230	Wei Zhang, Xin Jiang, Chunjing Xu, and Hang Xu.	1286
1231	2022. Wukong: A 100 million large-scale chinese	1287
1232	cross-modal pre-training benchmark . In <i>Advances in</i>	1288
1233	<i>Neural Information Processing Systems 35: Annual</i>	1289
1234	<i>Conference on Neural Information Processing Sys-</i>	1290
1235	<i>tems 2022, NeurIPS 2022, New Orleans, LA, USA,</i>	1291
1236	<i>November 28 - December 9, 2022.</i>	
1237	Anmol Gulati, James Qin, Chung-Cheng Chiu, Niki	
1238	Parmar, Yu Zhang, Jiahui Yu, Wei Han, Shibo Wang,	1292
1239	Zhengdong Zhang, Yonghui Wu, and Ruoming Pang.	1293
1240	2020. Conformer: Convolution-augmented trans-	1294
1241	former for speech recognition . In <i>21st Annual Con-</i>	1295
1242	<i>ference of the International Speech Communication</i>	1296
1243	<i>Association, Interspeech 2020, Virtual Event, Shang-</i>	1297
1244	<i>hai, China, October 25-29, 2020, pages 5036–5040.</i>	
1245	ISCA.	
1246	Suriya Gunasekar, Yi Zhang, Jyoti Aneja, Caio	
1247	César Teodoro Mendes, Allie Del Giorno, Sivakanth	1304
1248	Gopi, Mojan Javaheripi, Piero Kauffmann, Gustavo	1305
1249	de Rosa, Olli Saarikivi, Adil Salim, Shital Shah,	1306
1250	Harkirat Singh Behl, Xin Wang, Sébastien Bubeck,	1307
1251	Ronen Eldan, Adam Tauman Kalai, Yin Tat Lee, and	1308
1252	Yuanzhi Li. 2023. Textbooks are all you need . <i>CoRR</i> ,	1309
1253	abs/2306.11644.	1310
1254	Danna Gurari, Qing Li, Abigale J. Stangl, Anhong Guo,	
1255	Chi Lin, Kristen Grauman, Jiebo Luo, and Jeffrey P.	1311
1256	Bigham. 2018. Vizwiz grand challenge: Answering	1312
1257	visual questions from blind people . In <i>2018 IEEE</i>	1313
1258	<i>Conference on Computer Vision and Pattern Recog-</i>	
1259	<i>nition, CVPR 2018, Salt Lake City, UT, USA, June</i>	1314
1260	<i>18-22, 2018, pages 3608–3617. Computer Vision</i>	1315
1261	Foundation / IEEE Computer Society.	1316
		1317
1262	Jiaming Han, Kaixiong Gong, Yiyuan Zhang, Jiaqi	
1263	Wang, Kaipeng Zhang, Dahua Lin, Yu Qiao, Peng	1318
1264	Gao, and Xiangyu Yue. 2024a. Onellm: One frame-	1319
1265	work to align all modalities with language . In <i>IEEE/CVF</i>	1320
1266	<i>Conference on Computer Vision and Pat-</i>	1321
1267	<i>tern Recognition, CVPR 2024, Seattle, WA, USA,</i>	1322
1268	<i>June 16-22, 2024, pages 26574–26585. IEEE.</i>	1323
		1324
1269	Jiaming Han, Renrui Zhang, Wenqi Shao, Peng Gao,	
1270	Peng Xu, Han Xiao, Kaipeng Zhang, Chris Liu, Song	1325
1271	Wen, Ziyu Guo, Xudong Lu, Shuai Ren, Yafei Wen,	1326
1272	Xiaoxin Chen, Xiangyu Yue, Hongsheng Li, and	1327
1273	Yu Qiao. 2023. Imagebind-llm: Multi-modality in-	1328
1274	struction tuning . <i>CoRR</i> , abs/2309.03905.	1329
		1330
		1331
		1332
	Zeyu Han, Chao Gao, Jinyang Liu, Jeff Zhang, and	
	Sai Qian Zhang. 2024b. Parameter-efficient fine-	1275
	tuning for large models: A comprehensive survey .	1276
	<i>CoRR</i> , abs/2403.14608.	1277
		1278
	Yingqing He, Zhaoyang Liu, Jingye Chen, Zeyue Tian,	
	Hongyu Liu, Xiaowei Chi, Runtao Liu, Ruibin Yuan,	1279
	Yazhou Xing, Wenhai Wang, Jifeng Dai, Yong Zhang,	1280
	Wei Xue, Qifeng Liu, Yike Guo, and Qifeng Chen.	1281
	2024. Llms meet multimodal generation and editing:	1282
	A survey . <i>CoRR</i> , abs/2405.19334.	1283
		1284
	Yining Hong, Zishuo Zheng, Peihao Chen, Yian Wang,	
	Junyan Li, and Chuang Gan. 2024. Multiply: A	1285
	multisensory object-centric embodied large language	1286
	model in 3d world . In <i>IEEE/CVF Conference on</i>	1287
	<i>Computer Vision and Pattern Recognition, CVPR</i>	1288
	<i>2024, Seattle, WA, USA, June 16-22, 2024, pages</i>	1289
	<i>26396–26406. IEEE.</i>	1290
		1291
	Wei-Ning Hsu, Benjamin Bolte, Yao-Hung Hubert Tsai,	
	Kushal Lakhotia, Ruslan Salakhutdinov, and Abdel-	1292
	rahman Mohamed. 2021. Hubert: Self-supervised	1293
	speech representation learning by masked prediction	1294
	of hidden units . <i>IEEE ACM Trans. Audio Speech</i>	1295
	<i>Lang. Process.</i> , 29:3451–3460.	1296
		1297
	Edward J. Hu, Yelong Shen, Phillip Wallis, Zeyuan	
	Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and	1298
	Weizhu Chen. 2022. Lora: Low-rank adaptation of	1299
	large language models . In <i>The Tenth International</i>	1300
	<i>Conference on Learning Representations, ICLR 2022,</i>	1301
	<i>Virtual Event, April 25-29, 2022. OpenReview.net.</i>	1302
		1303
	Jiangyong Huang, Silong Yong, Xiaojian Ma, Xiongkun	
	Linghu, Puhao Li, Yan Wang, Qing Li, Song-Chun	1304
	Zhu, Baoxiong Jia, and Siyuan Huang. 2024. An	1305
	embodied generalist agent in 3d world . In <i>Forty-</i>	1306
	<i>first International Conference on Machine Learning,</i>	1307
	<i>ICML 2024, Vienna, Austria, July 21-27, 2024. Open-</i>	1308
	<i>Review.net.</i>	1309
		1310
	Jiaxing Huang and Jingyi Zhang. 2024. A survey	
	on evaluation of multimodal large language models .	1311
	<i>CoRR</i> , abs/2408.15769.	1312
		1313
	Xiaoshui Huang, Sheng Li, Wentao Qu, Tong He, Yi-	
	fan Zuo, and Wanli Ouyang. 2022. Frozen CLIP	1314
	model is an efficient point cloud backbone . <i>CoRR</i> ,	1315
	abs/2212.04098.	1316
		1317
	Drew A. Hudson and Christopher D. Manning. 2019.	
	GQA: A new dataset for real-world visual reason-	1318
	ing and compositional question answering . In <i>IEEE</i>	1319
	<i>Conference on Computer Vision and Pattern Recog-</i>	1320
	<i>nition, CVPR 2019, Long Beach, CA, USA, June 16-20,</i>	1321
	<i>2019, pages 6700–6709. Computer Vision Founda-</i>	1322
	<i>tion / IEEE.</i>	1323
		1324
	Albert Q. Jiang, Alexandre Sablayrolles, Arthur Men-	
	sch, Chris Bamford, Devendra Singh Chaplot, Diego	1325
	de Las Casas, Florian Bressand, Gianna Lengyel,	1326
	Guillaume Lample, Lucile Saulnier, Léo Ren-	1327
	ard Lavaud, Marie-Anne Lachaux, Pierre Stock,	1328
	Teven Le Scao, Thibaut Lavril, Thomas Wang, Timo-	1329
	thée Lacroix, and William El Sayed. 2023. Mistral	1330
	7b . <i>CoRR</i> , abs/2310.06825.	1331
		1332

1333	Albert Q. Jiang, Alexandre Sablayrolles, Antoine	Chris Dongjoo Kim, Byeongchang Kim, Hyunmin Lee,	1391
1334	Roux, Arthur Mensch, Blanche Savary, Chris Bam-	and Gunhee Kim. 2019. Audiocaps: Generating	1392
1335	ford, Devendra Singh Chaplot, Diego de Las Casas,	captions for audios in the wild . In <i>Proceedings of</i>	1393
1336	Emma Bou Hanna, Florian Bressand, Gianna	<i>the 2019 Conference of the North American Chap-</i>	1394
1337	Lengyel, Guillaume Bour, Guillaume Lample,	<i>ter of the Association for Computational Linguistics:</i>	1395
1338	L��lio Renard Lavaud, Lucile Saulnier, Marie-	<i>Human Language Technologies, NAACL-HLT 2019,</i>	1396
1339	Anne Lachaux, Pierre Stock, Sandeep Subramanian,	<i>Minneapolis, MN, USA, June 2-7, 2019, Volume 1</i>	1397
1340	Sophia Yang, Szymon Antoniak, Teven Le Scao,	<i>(Long and Short Papers)</i> , pages 119–132. Associa-	1398
1341	Th��ophile Gervet, Thibaut Lavril, Thomas Wang,	tion for Computational Linguistics.	1399
1342	Timoth��e Lacroix, and William El Sayed. 2024a.		
1343	Mixtral of experts . <i>CoRR</i> , abs/2401.04088.		
1344	Jianping Jiang, Weiye Xiao, Zhengyu Lin, Huaizhong	Taemin Kim, WOYEOL BAEK, and Heeseok Oh.	1400
1345	Zhang, Tianxiang Ren, Yang Gao, Zhiqian Lin, Zhon-	2024a. Efficient generative multimodal integration	1401
1346	gang Cai, Lei Yang, and Ziwei Liu. 2024b. SO-	(EGMI): Enabling audio generation from text-image	1402
1347	LAMI: social vision-language-action modeling for	pairs through alignment with large language models .	1403
1348	immersive interaction with 3d autonomous charac-	In <i>Audio Imagination: NeurIPS 2024 Workshop AI-</i>	1404
1349	ters . <i>CoRR</i> , abs/2412.00174.	<i>Driven Speech, Music, and Sound Generation</i> .	1405
1350	Yang Jin, Zhicheng Sun, Kun Xu, Kun Xu, Liwei Chen,	Yeonju Kim, Se Jin Park, and Yong Man Ro. 2024b.	1406
1351	Hao Jiang, Quzhe Huang, Chengru Song, Yuliang	Empathetic response in audio-visual conversations	1407
1352	Liu, Di Zhang, Yang Song, Kun Gai, and Yadong	using emotion preference optimization and mamba-	1408
1353	Mu. 2024. Video-lavit: Unified video-language pre-	compressor . <i>CoRR</i> , abs/2412.17572.	1409
1354	training with decoupled visual-motional tokenization .		
1355	In <i>Forty-first International Conference on Machine</i>	Qiuqiang Kong, Yin Cao, Turab Iqbal, Yuxuan Wang,	1410
1356	<i>Learning, ICML 2024, Vienna, Austria, July 21-27,</i>	Wenwu Wang, and Mark D. Plumbley. 2020. Panns:	1411
1357	<i>2024</i> . OpenReview.net.	Large-scale pretrained audio neural networks for au-	1412
		dio pattern recognition . <i>IEEE ACM Trans. Audio</i>	1413
		<i>Speech Lang. Process.</i> , 28:2880–2894.	1414
1358	Dimosthenis Karatzas, Llu��s Gomez-Bigorda, Angelos	Ben Koska and Moj��m��r Horv��th. 2024. Towards multi-	1415
1359	Nicolaou, Suman K. Ghosh, Andrew D. Bagdanov,	modal mastery: A 4.5b parameter truly multi-modal	1416
1360	Masakazu Iwamura, Jiri Matas, Luk��s Neumann, Vi-	small language model . <i>CoRR</i> , abs/2411.05903.	1417
1361	jay Ramaseshan Chandrasekhar, Shijian Lu, Faisal		
1362	Shafait, Seiichi Uchida, and Ernest Valveny. 2015.	Ranjay Krishna, Kenji Hata, Frederic Ren, Li Fei-Fei,	1418
1363	ICDAR 2015 competition on robust reading . In <i>13th</i>	and Juan Carlos Niebles. 2017a. Dense-captioning	1419
1364	<i>International Conference on Document Analysis and</i>	events in videos . In <i>IEEE International Conference</i>	1420
1365	<i>Recognition, ICDAR 2015, Nancy, France, August</i>	<i>on Computer Vision, ICCV 2017, Venice, Italy, Oc-</i>	1421
1366	<i>23-26, 2015</i> , pages 1156–1160. IEEE Computer So-	<i>ttober 22-29, 2017</i> , pages 706–715. IEEE Computer	1422
1367	cociety.	Society.	1423
1368	Dimosthenis Karatzas, Faisal Shafait, Seiichi Uchida,	Ranjay Krishna, Yuke Zhu, Oliver Groth, Justin John-	1424
1369	Masakazu Iwamura, Llu��s Gomez i Bigorda,	son, Kenji Hata, Joshua Kravitz, Stephanie Chen,	1425
1370	Sergi Robles Mestre, Joan Mas, David Fern��ndez	Yannis Kalantidis, Li-Jia Li, David A. Shamma,	1426
1371	Mota, Jon Almaz��n, and Llu��s-Pere de las Heras.	Michael S. Bernstein, and Li Fei-Fei. 2017b. Vi-	1427
1372	2013. ICDAR 2013 robust reading competition . In	sual genome: Connecting language and vision using	1428
1373	<i>12th International Conference on Document Analy-</i>	crowdsourced dense image annotations . <i>Int. J. Com-</i>	1429
1374	<i>sis and Recognition, ICDAR 2013, Washington, DC,</i>	<i>put. Vis.</i> , 123(1):32–73.	1430
1375	<i>USA, August 25-28, 2013</i> , pages 1484–1493. IEEE	Jinxiang Lai, Jie Zhang, Jun Liu, Jian Li, Xiaocheng	1431
1376	Computer Society.	Lu, and Song Guo. 2024. Spider: Any-to-many mul-	1432
		timodal LLM . <i>CoRR</i> , abs/2411.09439.	1433
1377	Justin Kerr, Huang Huang, Albert Wilcox, Ryan Hoque,	Hugo Lauren��on, Lucile Saulnier, L��o Tronchon, Stas	1434
1378	Jeffrey Ichnowski, Roberto Calandra, and Ken Gold-	Bekman, Amanpreet Singh, Anton Lozhkov, Thomas	1435
1379	berg. 2023. Self-supervised visuo-tactile pretraining	Wang, Siddharth Karamcheti, Alexander M. Rush,	1436
1380	to locate and follow garment features . In <i>Robotics:</i>	Douwe Kiela, Matthieu Cord, and Victor Sanh. 2023.	1437
1381	<i>Science and Systems XIX, Daegu, Republic of Korea,</i>	OBELICS: an open web-scale filtered dataset of inter-	1438
1382	<i>July 10-14, 2023</i> .	leaved image-text documents . In <i>Advances in Neural</i>	1439
1383	Douwe Kiela, Hamed Firooz, Aravind Mohan, Vedanuj	<i>Information Processing Systems 36: Annual Confer-</i>	1440
1384	Goswami, Amanpreet Singh, Pratik Ringshia, and	<i>ence on Neural Information Processing Systems 2023,</i>	1441
1385	Davide Testuggine. 2020. The hateful memes chal-	<i>NeurIPS 2023, New Orleans, LA, USA, December 10</i>	1442
1386	lenge: Detecting hate speech in multimodal memes .	<i>- 16, 2023</i> .	1443
1387	In <i>Advances in Neural Information Processing Sys-</i>	Dong-Guw Lee, Myung-Hwan Jeon, Younggun Cho,	1444
1388	<i>tems 33: Annual Conference on Neural Information</i>	and Ayoung Kim. 2023. Edge-guided multi-domain	1445
1389	<i>Processing Systems 2020, NeurIPS 2020, December</i>	rgb-to-tir image translation for training vision tasks	1446
1390	<i>6-12, 2020, virtual</i> .		

1447 with challenging labels. In *IEEE International Conference on Robotics and Automation, ICRA 2023, London, UK, May 29 - June 2, 2023*, pages 8291–8298. IEEE.

1448

1449

1450

1451 Sicong Leng, Yun Xing, Zesen Cheng, Yang Zhou, Hang Zhang, Xin Li, Deli Zhao, Shijian Lu, Chunyan Miao, and Lidong Bing. 2024. [The curse of multi-modalities: Evaluating hallucinations of large multimodal models across language, visual, and audio](#). *CoRR*, abs/2410.12787.

1452

1453

1454

1455

1456

1457 Guangyao Li, Yake Wei, Yapeng Tian, Chenliang Xu, Ji-Rong Wen, and Di Hu. 2022a. [Learning to answer questions in dynamic audio-visual scenarios](#). In *IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2022, New Orleans, LA, USA, June 18-24, 2022*, pages 19086–19096. IEEE.

1458

1459

1460

1461

1462

1463 Guangyao Li, Yake Wei, Yapeng Tian, Chenliang Xu, Ji-Rong Wen, and Di Hu. 2022b. [Learning to answer questions in dynamic audio-visual scenarios](#). In *IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2022, New Orleans, LA, USA, June 18-24, 2022*, pages 19086–19096. IEEE.

1464

1465

1466

1467

1468

1469 Jian Li and Weiheng Lu. 2024. [A survey on benchmarks of multimodal large language models](#). *CoRR*, abs/2408.08632.

1470

1471

1472 Junnan Li, Dongxu Li, Caiming Xiong, and Steven C. H. Hoi. 2022c. [BLIP: bootstrapping language-image pre-training for unified vision-language understanding and generation](#). In *International Conference on Machine Learning, ICML 2022, 17-23 July 2022, Baltimore, Maryland, USA*, volume 162 of *Proceedings of Machine Learning Research*, pages 12888–12900. PMLR.

1473

1474

1475

1476

1477

1478

1479

1480 Kunchang Li, Yali Wang, Yinan He, Yizhuo Li, Yi Wang, Yi Liu, Zun Wang, Jilan Xu, Guo Chen, Ping Lou, Limin Wang, and Yu Qiao. 2024a. [Mvbench: A comprehensive multi-modal video understanding benchmark](#). In *IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2024, Seattle, WA, USA, June 16-22, 2024*, pages 22195–22206. IEEE.

1481

1482

1483

1484

1485

1486

1487

1488 Linjie Li, Yen-Chun Chen, Yu Cheng, Zhe Gan, Licheng Yu, and Jingjing Liu. 2020. [HERO: hierarchical encoder for video+language omni-representation pre-training](#). In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing, EMNLP 2020, Online, November 16-20, 2020*, pages 2046–2065. Association for Computational Linguistics.

1489

1490

1491

1492

1493

1494

1495

1496 Mingsheng Li, Xin Chen, Chi Zhang, Sijin Chen, Hongyuan Zhu, Fukun Yin, Zhuoyuan Li, Gang Yu, and Tao Chen. 2024b. [M3dbench: Towards omni 3d assistant with interleaved multi-modal instructions](#). In *Computer Vision - ECCV 2024 - 18th European Conference, Milan, Italy, September 29-October 4, 2024, Proceedings, Part LVIII*, volume 15116 of *Lecture Notes in Computer Science*, pages 41–59. Springer.

1497

1498

1499

1500

1501

1502

1503

1504

Yadong Li, Jun Liu, Tao Zhang, Tao Zhang, Song Chen, Tianpeng Li, Zehuan Li, Lijun Liu, Lingfeng Ming, Guosheng Dong, Da Pan, Chong Li, Yuanbo Fang, Dongdong Kuang, Mingrui Wang, Chenglin Zhu, Youwei Zhang, Hongyu Guo, Fengyu Zhang, Yuran Wang, Bowen Ding, Wei Song, Xu Li, Yuqi Huo, Zheng Liang, Shusen Zhang, Xin Wu, Shuai Zhao, Linchu Xiong, Yozhen Wu, Jiahui Ye, Wenhao Lu, Bowen Li, Yan Zhang, Yaqi Zhou, Xin Chen, Lei Su, Hongda Zhang, Fuzhong Chen, Xuezhen Dong, Na Nie, Zhiying Wu, Bin Xiao, Ting Li, Shunya Dang, Ping Zhang, Yijia Sun, Jincheng Wu, Jinjie Yang, Xionghai Lin, Zhi Ma, Kegeng Wu, Jiali, Aiyuan Yang, Hui Liu, Jianqiang Zhang, Xiaoxi Chen, Guangwei Ai, Wentao Zhang, Yicong Chen, Xiaoqin Huang, Kun Li, Wenjing Luo, Yifei Duan, Lingling Zhu, Ran Xiao, Zhe Su, Jiani Pu, Dian Wang, Xu Jia, Tianyu Zhang, Mengyu Ai, Mang Wang, Yujing Qiao, Lei Zhang, Yanjun Shen, Fan Yang, Miao Zhen, Yijie Zhou, Mingyang Chen, Fei Li, Chenzheng Zhu, Keer Lu, Yaqi Zhao, Hao Liang, Youquan Li, Yanzhao Qin, Linzhuang Sun, Jianhua Xu, Haoze Sun, Mingan Lin, Zenan Zhou, and Weipeng Chen. 2025. [Baichuan-omni-1.5 technical report](#). *Preprint*, arXiv:2501.15368.

1505

1506

1507

1508

1509

1510

1511

1512

1513

1514

1515

1516

1517

1518

1519

1520

1521

1522

1523

1524

1525

1526

1527

1528

1529

Yadong Li, Haoze Sun, Mingan Lin, Tianpeng Li, Guosheng Dong, Tao Zhang, Bowen Ding, Wei Song, Zhenglin Cheng, Yuqi Huo, Song Chen, Xu Li, Da Pan, Shusen Zhang, Xin Wu, Zheng Liang, Jun Liu, Tao Zhang, Keer Lu, Yaqi Zhao, Yanjun Shen, Fan Yang, Kaicheng Yu, Tao Lin, Jianhua Xu, Zenan Zhou, and Weipeng Chen. 2024c. [Baichuan-omni technical report](#). *CoRR*, abs/2410.08565.

1530

1531

1532

1533

1534

1535

1536

1537

Yifan Li, Yifan Du, Kun Zhou, Jinpeng Wang, Wayne Xin Zhao, and Ji-Rong Wen. 2023. [Evaluating object hallucination in large vision-language models](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing, EMNLP 2023, Singapore, December 6-10, 2023*, pages 292–305. Association for Computational Linguistics.

1538

1539

1540

1541

1542

1543

1544

1545

Yinghao Aaron Li, Cong Han, and Nima Mesgarani. 2022d. [Styletts: A style-based generative model for natural and diverse text-to-speech synthesis](#). *CoRR*, abs/2205.15439.

1546

1547

1548

1549

Yizhi Li, Ruibin Yuan, Ge Zhang, Yinghao Ma, Xingran Chen, Hanzhi Yin, Chenghao Xiao, Chenghua Lin, Anton Ragni, Emmanouil Benetos, Norbert Gyenge, Roger B. Dannenberg, Ruibo Liu, Wenhao Chen, Gus Xia, Yemin Shi, Wenhao Huang, Zili Wang, Yike Guo, and Jie Fu. 2024d. [MERT: acoustic music understanding model with large-scale self-supervised training](#). In *The Twelfth International Conference on Learning Representations, ICLR 2024, Vienna, Austria, May 7-11, 2024*. OpenReview.net.

1550

1551

1552

1553

1554

1555

1556

1557

1558

1559

Yizhi Li, Ge Zhang, Yinghao Ma, Ruibin Yuan, Kang Zhu, Hangyu Guo, Yiming Liang, Jiaheng Liu, Jian Yang, Siwei Wu, Xingwei Qu, Jinjie Shi, Xinyue Zhang, Zhenzhu Yang, Xiangzhou Wang, Zhaoxiang Zhang, Zachary Liu, Emmanouil Benetos, Wenhao

1560

1561

1562

1563

1564

1565	Huang, and Chenghua Lin. 2024e. Omnibench: Towards the future of universal omni-language models . <i>CoRR</i> , abs/2409.15272.	<i>Proceedings of Machine Learning Research</i> , pages 21450–21474. PMLR.	1622
1566			1623
1567			
1568	Yunxin Li, Shenyuan Jiang, Baotian Hu, Longyue Wang, Wanqi Zhong, Wenhan Luo, Lin Ma, and Min Zhang. 2024f. Uni-moe: Scaling unified multimodal llms with mixture of experts . <i>CoRR</i> , abs/2405.11273.	Haotian Liu, Chunyuan Li, Yuheng Li, and Yong Jae Lee. 2024a. Improved baselines with visual instruction tuning . In <i>IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2024, Seattle, WA, USA, June 16-22, 2024</i> , pages 26286–26296. IEEE.	1624
1569			1625
1570			1626
1571			1627
1572	Zhaowei Li, Wei Wang, Yiqing Cai, Qi Xu, Pengyu Wang, Dong Zhang, Hang Song, Botian Jiang, Zhida Huang, and Tao Wang. 2024g. Unifiedm-llm: Enabling unified representation for multi-modal multi-tasks with large language model . <i>CoRR</i> , abs/2408.02503.		1628
1573			1629
1574		Haotian Liu, Chunyuan Li, Qingyang Wu, and Yong Jae Lee. 2023c. Visual instruction tuning . In <i>Advances in Neural Information Processing Systems 36: Annual Conference on Neural Information Processing Systems 2023, NeurIPS 2023, New Orleans, LA, USA, December 10 - 16, 2023</i> .	1630
1575			1631
1576			1632
1577			1633
1578	Zhaowei Li, Qi Xu, Dong Zhang, Hang Song, Yiqing Cai, Qi Qi, Ran Zhou, Junting Pan, Zefeng Li, Vu Tu, Zhida Huang, and Tao Wang. 2024h. Groundinggpt: Language enhanced multi-modal grounding model . In <i>Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers), ACL 2024, Bangkok, Thailand, August 11-16, 2024</i> , pages 6657–6678. Association for Computational Linguistics.		1634
1579			1635
1580		Shansong Liu, Atin Sakkeer Hussain, Qilong Wu, Chen-shuo Sun, and Ying Shan. 2024b. Mumu-llama: Multi-modal music understanding and generation via large language models . <i>CoRR</i> , abs/2412.06660.	1636
1581			1637
1582			1638
1583			1639
1584		Yaofang Liu, Xiaodong Cun, Xuebo Liu, Xintao Wang, Yong Zhang, Haoxin Chen, Yang Liu, Tiejong Zeng, Raymond Chan, and Ying Shan. 2024c. Evalcrafter: Benchmarking and evaluating large video generation models . In <i>IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2024, Seattle, WA, USA, June 16-22, 2024</i> , pages 22139–22149. IEEE.	1640
1585			1641
1586			1642
1587	Tian Liang, Jing Huang, Ming Kong, Luyuan Chen, and Qiang Zhu. 2024. Querying as prompt: Parameter-efficient learning for multimodal language model . In <i>IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2024, Seattle, WA, USA, June 16-22, 2024</i> , pages 26845–26855. IEEE.		1643
1588			1644
1589			1645
1590			1646
1591			1647
1592			
1593	Ji Lin, Hongxu Yin, Wei Ping, Pavlo Molchanov, Mohammad Shoeybi, and Song Han. 2024. VILA: on pre-training for visual language models . In <i>IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2024, Seattle, WA, USA, June 16-22, 2024</i> , pages 26679–26689. IEEE.	Yuan Liu, Haodong Duan, Yuanhan Zhang, Bo Li, Songyang Zhang, Wangbo Zhao, Yike Yuan, Jiaqi Wang, Conghui He, Ziwei Liu, Kai Chen, and Dahua Lin. 2024d. Mmbench: Is your multi-modal model an all-around player? In <i>Computer Vision - ECCV 2024 - 18th European Conference, Milan, Italy, September 29-October 4, 2024, Proceedings, Part VI</i> , volume 15064 of <i>Lecture Notes in Computer Science</i> , pages 216–233. Springer.	1648
1594			1649
1595			1650
1596			1651
1597			1652
1598			1653
1599	Tsung-Yi Lin, Michael Maire, Serge J. Belongie, James Hays, Pietro Perona, Deva Ramanan, Piotr Dollár, and C. Lawrence Zitnick. 2014. Microsoft COCO: common objects in context . In <i>Computer Vision - ECCV 2014 - 13th European Conference, Zurich, Switzerland, September 6-12, 2014, Proceedings, Part V</i> , volume 8693 of <i>Lecture Notes in Computer Science</i> , pages 740–755. Springer.		1654
1600			1655
1601			1656
1602		Ze Liu, Jia Ning, Yue Cao, Yixuan Wei, Zheng Zhang, Stephen Lin, and Han Hu. 2022a. Video swin transformer . In <i>IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2022, New Orleans, LA, USA, June 18-24, 2022</i> , pages 3192–3201. IEEE.	1657
1603			1658
1604			1659
1605			1660
1606			1661
1607			1662
1608	Samuel Lipping, Parthasaarathy Sudarsanam, Konstantinos Drossos, and Tuomas Virtanen. 2022. Clotho-aqa: A crowdsourced dataset for audio question answering . In <i>30th European Signal Processing Conference, EUSIPCO 2022, Belgrade, Serbia, August 29 - Sept. 2, 2022</i> , pages 1140–1144. IEEE.		1663
1609			1664
1610			1665
1611			1666
1612			1667
1613	Fangyu Liu, Guy Emerson, and Nigel Collier. 2023a. Visual spatial reasoning . <i>Trans. Assoc. Comput. Linguistics</i> , 11:635–651.		1668
1614			
1615			
1616	Haohe Liu, Zehua Chen, Yi Yuan, Xinhao Mei, Xubo Liu, Danilo P. Mandic, Wenwu Wang, and Mark D. Plumbley. 2023b. Audioldm: Text-to-audio generation with latent diffusion models . In <i>International Conference on Machine Learning, ICML 2023, 23-29 July 2023, Honolulu, Hawaii, USA</i> , volume 202 of	Jiasen Lu, Christopher Clark, Sangho Lee, Zichen Zhang, Savya Khosla, Ryan Marten, Derek Hoiem, and Aniruddha Kembhavi. 2024a. Unified-io 2: Scaling autoregressive multimodal models with vision, language, audio, and action . In <i>IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2024, Seattle, WA, USA, June 16-22, 2024</i> , pages 26429–26445. IEEE.	1669
1617			1670
1618			1671
1619			1672
1620			1673
1621			1674
			1675
			1676

- Pan Lu, Swaroop Mishra, Tanglin Xia, Liang Qiu, Kai-Wei Chang, Song-Chun Zhu, Oyvind Tafjord, Peter Clark, and Ashwin Kalyan. 2022. [Learn to explain: Multimodal reasoning via thought chains for science question answering](#). In *Advances in Neural Information Processing Systems 35: Annual Conference on Neural Information Processing Systems 2022, NeurIPS 2022, New Orleans, LA, USA, November 28 - December 9, 2022*.
- Pan Lu, Liang Qiu, Jiaqi Chen, Tanglin Xia, Yizhou Zhao, Wei Zhang, Zhou Yu, Xiaodan Liang, and Song-Chun Zhu. 2021. [Iconqa: A new benchmark for abstract diagram understanding and visual language reasoning](#). In *Proceedings of the Neural Information Processing Systems Track on Datasets and Benchmarks 1, NeurIPS Datasets and Benchmarks 2021, December 2021, virtual*.
- Yichen Lu, Jiaqi Song, Xuankai Chang, Hengwei Bian, Soumi Maiti, and Shinji Watanabe. 2024b. [Syneslm: A unified approach for audio-visual speech recognition and translation via language model and synthetic data](#). *CoRR*, abs/2408.00624.
- Mingshuang Luo, Ruibing Hou, Hong Chang, Zimo Liu, Yaowei Wang, and Shiguang Shan. 2024. [M³gpt: An advanced multimodal, multitask framework for motion comprehension and generation](#). *CoRR*, abs/2405.16273.
- Ruipu Luo, Ziwang Zhao, Min Yang, Junwei Dong, Minghui Qiu, Pengcheng Lu, Tao Wang, and Zhongyu Wei. 2023a. [Valley: Video assistant with large language model enhanced ability](#). *CoRR*, abs/2306.07207.
- Run Luo, Ting-En Lin, Haonan Zhang, Yuchuan Wu, Xiong Liu, Min Yang, Yongbin Li, Longze Chen, Jiaming Li, Lei Zhang, Yangyi Chen, Hamid Alinejad-Rokny, and Fei Huang. 2025. [Openomni: Large language models pivot zero-shot omnimodal alignment across language with real-time self-aware emotional speech synthesis](#). *Preprint*, arXiv:2501.04561.
- Tiange Luo, Chris Rockwell, Honglak Lee, and Justin Johnson. 2023b. [Scalable 3d captioning with pre-trained models](#). In *Advances in Neural Information Processing Systems 36: Annual Conference on Neural Information Processing Systems 2023, NeurIPS 2023, New Orleans, LA, USA, December 10 - 16, 2023*.
- Zhengxiong Luo, Dayou Chen, Yingya Zhang, Yan Huang, Liang Wang, Yujun Shen, Deli Zhao, Jingren Zhou, and Tieniu Tan. 2023c. [Videofusion: Decomposed diffusion models for high-quality video generation](#). *CoRR*, abs/2303.08320.
- Chenyang Lyu, Minghao Wu, Longyue Wang, Xinting Huang, Bingshuai Liu, Zefeng Du, Shuming Shi, and Zhaopeng Tu. 2023. [Macaw-llm: Multi-modal language modeling with image, audio, video, and text integration](#). *CoRR*, abs/2306.09093.
- Xianzheng Ma, Yash Bhalgat, Brandon Smart, Shuai Chen, Xinghui Li, Jian Ding, Jindong Gu, Dave Zhenyu Chen, Songyou Peng, Jia-Wang Bian, Philip H. S. Torr, Marc Pollefeys, Matthias Nießner, Ian D. Reid, Angel X. Chang, Iro Laina, and Victor Adrian Prisacariu. 2024. [When llms step into the 3d world: A survey and meta-analysis of 3d tasks via multi-modal large language models](#). *CoRR*, abs/2405.10255.
- Xiaojian Ma, Silong Yong, Zilong Zheng, Qing Li, Yitao Liang, Song-Chun Zhu, and Siyuan Huang. 2023. [SQA3D: situated question answering in 3d scenes](#). In *The Eleventh International Conference on Learning Representations, ICLR 2023, Kigali, Rwanda, May 1-5, 2023*. OpenReview.net.
- Muhammad Maaz, Hanoona Abdul Rasheed, Salman Khan, and Fahad Khan. 2024. [Video-chatgpt: Towards detailed video understanding via large vision and language models](#). In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers), ACL 2024, Bangkok, Thailand, August 11-16, 2024*, pages 12585–12602. Association for Computational Linguistics.
- Karttikeya Mangalam, Raiymbek Akshulakov, and Jitendra Malik. 2023. [Egoschema: A diagnostic benchmark for very long-form video language understanding](#). In *Advances in Neural Information Processing Systems 36: Annual Conference on Neural Information Processing Systems 2023, NeurIPS 2023, New Orleans, LA, USA, December 10 - 16, 2023*.
- Junhua Mao, Jonathan Huang, Alexander Toshev, Oana Camburu, Alan L. Yuille, and Kevin Murphy. 2016. [Generation and comprehension of unambiguous object descriptions](#). In *2016 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2016, Las Vegas, NV, USA, June 27-30, 2016*, pages 11–20. IEEE Computer Society.
- Kenneth Marino, Mohammad Rastegari, Ali Farhadi, and Roozbeh Mottaghi. 2019. [OK-VQA: A visual question answering benchmark requiring external knowledge](#). In *IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2019, Long Beach, CA, USA, June 16-20, 2019*, pages 3195–3204. Computer Vision Foundation / IEEE.
- Minesh Mathew, Dimosthenis Karatzas, and C. V. Jawahar. 2021. [Docvqa: A dataset for VQA on document images](#). In *IEEE Winter Conference on Applications of Computer Vision, WACV 2021, Waikoloa, HI, USA, January 3-8, 2021*, pages 2199–2208. IEEE.
- Xinhao Mei, Chutong Meng, Haohe Liu, Qiuqiang Kong, Tom Ko, Chengqi Zhao, Mark D. Plumbley, Yuexian Zou, and Wenwu Wang. 2024. [Wavcaps: A chatgpt-assisted weakly-labelled audio captioning dataset for audio-language multimodal research](#). *IEEE ACM Trans. Audio Speech Lang. Process.*, 32:3339–3354.

- Annamaria Mesaros, Toni Heittola, and Tuomas Virtanen. 2016. [TUT database for acoustic scene classification and sound event detection](#). In *24th European Signal Processing Conference, EUSIPCO 2016, Budapest, Hungary, August 29 - September 2, 2016*, pages 1128–1132. IEEE.
- OpenBMB MiniCPM-o Team. 2025. Minicpm-o 2.6: A gpt-4o level mllm for vision, speech, and multimodal live streaming on your phone. <https://github.com/OpenBMB/MiniCPM-o>. Accessed: 2025-02-10.
- Anand Mishra, Karteek Alahari, and C. V. Jawahar. 2012. [Scene text recognition using higher order language priors](#). In *British Machine Vision Conference, BMVC 2012, Surrey, UK, September 3-7, 2012*, pages 1–11. BMVA Press.
- Anand Mishra, Shashank Shekhar, Ajeet Kumar Singh, and Anirban Chakraborty. 2019. [OCR-VQA: visual question answering by reading text in images](#). In *2019 International Conference on Document Analysis and Recognition, ICDAR 2019, Sydney, Australia, September 20-25, 2019*, pages 947–952. IEEE.
- Seungwhan Moon, Andrea Madotto, Zhaojiang Lin, Alireza Dirafzoon, Aparajita Saraf, Amy Bearman, and Babak Damavandi. 2022. [IMU2CLIP: multimodal contrastive learning for IMU motion sensors from egocentric videos and text](#). *CoRR*, abs/2210.14395.
- Seungwhan Moon, Andrea Madotto, Zhaojiang Lin, Tushar Nagarajan, Matt Smith, Shashank Jain, Chun-Fu Yeh, Prakash Murugesan, Peyman Heidari, Yue Liu, Kavya Srinet, Babak Damavandi, and Anuj Kumar. 2024. [Anymal: An efficient and scalable any-modality augmented language model](#). In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing: EMNLP 2024 - Industry Track, Miami, Florida, USA, November 12-16, 2024*, pages 1314–1332. Association for Computational Linguistics.
- OpenAI. Hello gpt-4o. <https://openai.com/index/hello-gpt-4o/>. Accessed: January 3, 2025.
- OpenAI. 2023a. Chatgpt. Technical report, OpenAI.
- OpenAI. 2023b. [GPT-4 technical report](#). *CoRR*, abs/2303.08774.
- OpenGV. 2024. [Interomni](#). Accessed: 2024-07-27.
- Maxime Oquab, Timothée Darcet, Théo Moutakanni, Huy V. Vo, Marc Szafraniec, Vasil Khalidov, Pierre Fernandez, Daniel Haziza, Francisco Massa, Alaaeldin El-Nouby, Mido Assran, Nicolas Ballas, Wojciech Galuba, Russell Howes, Po-Yao Huang, Shang-Wen Li, Ishan Misra, Michael Rabbat, Vasu Sharma, Gabriel Synnaeve, Hu Xu, Hervé Jégou, Julien Mairal, Patrick Labatut, Armand Joulin, and Piotr Bojanowski. 2024. [Dinov2: Learning robust visual features without supervision](#). *Trans. Mach. Learn. Res.*, 2024.
- Vicente Ordonez, Girish Kulkarni, and Tamara L. Berg. 2011. [Im2text: Describing images using 1 million captioned photographs](#). In *Advances in Neural Information Processing Systems 24: 25th Annual Conference on Neural Information Processing Systems 2011. Proceedings of a meeting held 12-14 December 2011, Granada, Spain*, pages 1143–1151.
- Aitor Ormazabal, Che Zheng, Cyprien de Masson d’Autume, Dani Yogatama, Deyu Fu, Donovan Ong, Eric Chen, Eugenie Lamprecht, Hai Pham, Isaac Ong, Kaloyan Aleksiev, Lei Li, Matthew Henderson, Max Bain, Mikel Artetxe, Nishant Relan, Piotr Padlewski, Qi Liu, Ren Chen, Samuel Phua, Yazheng Yang, Yi Tay, Yuqi Wang, Zhongkai Zhu, and Zhihui Xie. 2024. [Reka core, flash, and edge: A series of powerful multimodal language models](#). *CoRR*, abs/2404.12387.
- Artemis Panagopoulou, Le Xue, Ning Yu, Junnan Li, Dongxu Li, Shafiq Joty, Ran Xu, Silvio Savarese, Caiming Xiong, and Juan Carlos Niebles. 2024. [X-instructblip: A framework for aligning image, 3d, audio, video to llms and its emergent cross-modal reasoning](#). In *Computer Vision - ECCV 2024 - 18th European Conference, Milan, Italy, September 29-October 4, 2024, Proceedings, Part XLV*, volume 15103 of *Lecture Notes in Computer Science*, pages 177–197. Springer.
- Vassil Panayotov, Guoguo Chen, Daniel Povey, and Sanjeev Khudanpur. 2015. [Librispeech: An ASR corpus based on public domain audio books](#). In *2015 IEEE International Conference on Acoustics, Speech and Signal Processing, ICASSP 2015, South Brisbane, Queensland, Australia, April 19-24, 2015*, pages 5206–5210. IEEE.
- Haozhou Pang, Tianwei Ding, Lanshan He, Ming Tao, Lu Zhang, and Qi Gan. 2024. [LLM gesticulator: Leveraging large language models for scalable and controllable co-speech gesture synthesis](#). *CoRR*, abs/2410.10851.
- Viorica Patraucean, Lucas Smaira, Ankush Gupta, Adrià Recasens, Larisa Markeeva, Dylan Banarse, Skanda Koppula, Joseph Heyward, Mateusz Malinowski, Yi Yang, Carl Doersch, Tatiana Matejovicova, Yury Sulsky, Antoine Miech, Alexandre Fréchet, Hanna Klimczak, Raphael Koster, Junlin Zhang, Stephanie Winkler, Yusuf Aytar, Simon Osindero, Dima Damen, Andrew Zisserman, and João Carreira. 2023. [Perception test: A diagnostic benchmark for multimodal video models](#). In *Advances in Neural Information Processing Systems 36: Annual Conference on Neural Information Processing Systems 2023, NeurIPS 2023, New Orleans, LA, USA, December 10 - 16, 2023*.
- Zhiliang Peng, Li Dong, Hangbo Bao, Qixiang Ye, and Furu Wei. 2022. [Beit v2: Masked image modeling with vector-quantized visual tokenizers](#). *CoRR*, abs/2208.06366.
- Federico Perazzi, Jordi Pont-Tuset, Brian McWilliams, Luc Van Gool, Markus H. Gross, and Alexander

1904	Sorkine-Hornung. 2016. A benchmark dataset and evaluation methodology for video object segmentation . In <i>2016 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2016, Las Vegas, NV, USA, June 27-30, 2016</i> , pages 724–732. IEEE Computer Society.	Colin Raffel, Noam Shazeer, Adam Roberts, Katherine Lee, Sharan Narang, Michael Matena, Yanqi Zhou, Wei Li, and Peter J. Liu. 2020. Exploring the limits of transfer learning with a unified text-to-text transformer . <i>J. Mach. Learn. Res.</i> , 21:140:1–140:67.	1961 1962 1963 1964 1965
1910	Trung Quy Phan, Palaiahnakote Shivakumara, Shangxuan Tian, and Chew Lim Tan. 2013. Recognizing text with perspective distortion in natural scenes . In <i>IEEE International Conference on Computer Vision, ICCV 2013, Sydney, Australia, December 1-8, 2013</i> , pages 569–576. IEEE Computer Society.	Aditya Ramesh, Mikhail Pavlov, Gabriel Goh, Scott Gray, Chelsea Voss, Alec Radford, Mark Chen, and Ilya Sutskever. 2021. Zero-shot text-to-image generation . <i>CoRR</i> , abs/2102.12092.	1966 1967 1968 1969
1916	Karol J. Piczak. 2015. ESC: dataset for environmental sound classification . In <i>Proceedings of the 23rd Annual ACM Conference on Multimedia Conference, MM '15, Brisbane, Australia, October 26 - 30, 2015</i> , pages 1015–1018. ACM.	René Ranftl, Alexey Bochkovskiy, and Vladlen Koltun. 2021. Vision transformers for dense prediction . In <i>2021 IEEE/CVF International Conference on Computer Vision, ICCV 2021, Montreal, QC, Canada, October 10-17, 2021</i> , pages 12159–12168. IEEE.	1970 1971 1972 1973 1974
1921	Bryan A. Plummer, Liwei Wang, Chris M. Cervantes, Juan C. Caicedo, Julia Hockenmaier, and Svetlana Lazebnik. 2015. Flickr30k entities: Collecting region-to-phrase correspondences for richer image-to-sentence models . In <i>2015 IEEE International Conference on Computer Vision, ICCV 2015, Santiago, Chile, December 7-13, 2015</i> , pages 2641–2649. IEEE Computer Society.	Machel Reid, Nikolay Savinov, Denis Teplyashin, Dmitry Lepikhin, Timothy P. Lillicrap, Jean-Baptiste Alayrac, Radu Soricut, Angeliki Lazaridou, Orhan Firat, Julian Schrittwieser, Ioannis Antonoglou, Rohan Anil, Sebastian Borgeaud, Andrew M. Dai, Katie Millican, Ethan Dyer, Mia Glaese, Thibault Sottiaux, Benjamin Lee, Fabio Viola, Malcolm Reynolds, Yuanzhong Xu, James Molloy, Jilin Chen, Michael Isard, Paul Barham, Tom Hennigan, Ross McIlroy, Melvin Johnson, Johan Schalkwyk, Eli Collins, Eliza Rutherford, Erica Moreira, Kareem Ayoub, Megha Goel, Clemens Meyer, Gregory Thornton, Zhen Yang, Henryk Michalewski, Zaheer Abbas, Nathan Schucher, Ankesh Anand, Richard Ives, James Keeling, Karel Lenc, Salem Haykal, Siamak Shakeri, Pranav Shyam, Aakanksha Chowdhery, Roman Ring, Stephen Spencer, Eren Sezener, and et al. 2024. Gemini 1.5: Unlocking multimodal understanding across millions of tokens of context . <i>CoRR</i> , abs/2403.05530.	1975 1976 1977 1978 1979 1980 1981 1982 1983 1984 1985 1986 1987 1988 1989 1990 1991 1992 1993 1994
1929	Vineel Pratap, Qiantong Xu, Anuroop Sriram, Gabriel Synnaeve, and Ronan Collobert. 2020. MLS: A large-scale multilingual dataset for speech research . In <i>21st Annual Conference of the International Speech Communication Association, Interspeech 2020, Virtual Event, Shanghai, China, October 25-29, 2020</i> , pages 2757–2761. ISCA.	Yong Ren, Tao Wang, Jiangyan Yi, Le Xu, Jianhua Tao, Chu Yuan Zhang, and Junzuo Zhou. 2024. Fewer-token neural speech codec with time-invariant codes . In <i>IEEE International Conference on Acoustics, Speech and Signal Processing, ICASSP 2024, Seoul, Republic of Korea, April 14-19, 2024</i> , pages 12737–12741. IEEE.	1995 1996 1997 1998 1999 2000 2001
1936	Filip Radenovic, Abhimanyu Dubey, Abhishek Kadian, Todor Mihaylov, Simon Vandenhende, Yash Patel, Yi Wen, Vignesh Ramanathan, and Dhruv Mahajan. 2023. Filtering, distillation, and hard negatives for vision-language pre-training . In <i>IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2023, Vancouver, BC, Canada, June 17-24, 2023</i> , pages 6967–6977. IEEE.	Anhar Risnumawan, Palaiahnakote Shivakumara, Chee Seng Chan, and Chew Lim Tan. 2014. A robust arbitrary text detection system for natural scene images . <i>Expert Syst. Appl.</i> , 41(18):8027–8048.	2002 2003 2004 2005
1944	Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, Gretchen Krueger, and Ilya Sutskever. 2021. Learning transferable visual models from natural language supervision . In <i>Proceedings of the 38th International Conference on Machine Learning, ICML 2021, 18-24 July 2021, Virtual Event</i> , volume 139 of <i>Proceedings of Machine Learning Research</i> , pages 8748–8763. PMLR.	Robin Rombach, Andreas Blattmann, Dominik Lorenz, Patrick Esser, and Björn Ommer. 2022. High-resolution image synthesis with latent diffusion models . In <i>IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2022, New Orleans, LA, USA, June 18-24, 2022</i> , pages 10674–10685. IEEE.	2006 2007 2008 2009 2010 2011 2012
1954	Alec Radford, Jong Wook Kim, Tao Xu, Greg Brockman, Christine McLeavey, and Ilya Sutskever. 2023. Robust speech recognition via large-scale weak supervision . In <i>International Conference on Machine Learning, ICML 2023, 23-29 July 2023, Honolulu, Hawaii, USA</i> , volume 202 of <i>Proceedings of Machine Learning Research</i> , pages 28492–28518. PMLR.	Nataniel Ruiz, Yuanzhen Li, Varun Jampani, Yael Pritch, Michael Rubinstein, and Kfir Aberman. 2023. Dreambooth: Fine tuning text-to-image diffusion models for subject-driven generation . In <i>IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2023, Vancouver, BC, Canada, June 17-24, 2023</i> , pages 22500–22510. IEEE.	2013 2014 2015 2016 2017 2018 2019

2020	RVC-Boss. Gpt-sovits. https://github.com/RVC-Boss/GPT-SoVITS . Accessed: January 3, 2025.	Zhan Shi, Xu Zhou, Xipeng Qiu, and Xiaodan Zhu. 2020. Improving image captioning with better use of captions . <i>CoRR</i> , abs/2006.11807.	2075 2076 2077
2023	Ramon Sanabria, Ozan Caglayan, Shruti Palaskar, Desmond Elliott, Loïc Barrault, Lucia Specia, and Florian Metze. 2018. How2: A large-scale dataset for multimodal language understanding . <i>CoRR</i> , abs/1811.00347.	Fangxun Shu, Lei Zhang, Hao Jiang, and Cihang Xie. 2023. Audio-visual LLM for video understanding . <i>CoRR</i> , abs/2312.06720.	2078 2079 2080
2026		Mustafa Shukor, Corentin Dancette, and Matthieu Cord. 2023. ep-alm: Efficient perceptual augmentation of language models . In <i>IEEE/CVF International Conference on Computer Vision, ICCV 2023, Paris, France, October 1-6, 2023</i> , pages 21999–22012. IEEE.	2081 2082 2083 2084 2085 2086
2028	Christoph Schuhmann, Romain Beaumont, Richard Vencu, Cade Gordon, Ross Wightman, Mehdi Cherti, Theo Coombes, Aarush Katta, Clayton Mullis, Mitchell Wortsman, Patrick Schramowski, Srivatsa Kundurthy, Katherine Crowson, Ludwig Schmidt, Robert Kaczmarczyk, and Jenia Jitsev. 2022. LAION-5B: an open large-scale dataset for training next generation image-text models . In <i>Advances in Neural Information Processing Systems 35: Annual Conference on Neural Information Processing Systems 2022, NeurIPS 2022, New Orleans, LA, USA, November 28 - December 9, 2022</i> .	Gunnar A. Sigurdsson, Gül Varol, Xiaolong Wang, Ali Farhadi, Ivan Laptev, and Abhinav Gupta. 2016. Hollywood in homes: Crowdsourcing data collection for activity understanding . In <i>Computer Vision - ECCV 2016 - 14th European Conference, Amsterdam, The Netherlands, October 11-14, 2016, Proceedings, Part I</i> , volume 9905 of <i>Lecture Notes in Computer Science</i> , pages 510–526. Springer.	2087 2088 2089 2090 2091 2092 2093 2094
2040	Christoph Schuhmann, Andreas Köpf, Richard Vencu, Theo Coombes, and Romain Beaumont. 2022b(a). Laion-aesthetics .	Nathan Silberman, Derek Hoiem, Pushmeet Kohli, and Rob Fergus. 2012. Indoor segmentation and support inference from RGBD images . In <i>Computer Vision - ECCV 2012 - 12th European Conference on Computer Vision, Florence, Italy, October 7-13, 2012, Proceedings, Part V</i> , volume 7576 of <i>Lecture Notes in Computer Science</i> , pages 746–760. Springer.	2095 2096 2097 2098 2099 2100 2101
2043	Christoph Schuhmann, Andreas Köpf, Richard Vencu, Theo Coombes, and Romain Beaumont. 2022b(b). Laion coco: 600m synthetic captions from laion2b-en .	Amanpreet Singh, Vivek Natarajan, Meet Shah, Yu Jiang, Xinlei Chen, Dhruv Batra, Devi Parikh, and Marcus Rohrbach. 2019. Towards VQA models that can read . In <i>IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2019, Long Beach, CA, USA, June 16-20, 2019</i> , pages 8317–8326. Computer Vision Foundation / IEEE.	2102 2103 2104 2105 2106 2107 2108
2047	Christoph Schuhmann, Richard Vencu, Romain Beaumont, Robert Kaczmarczyk, Clayton Mullis, Aarush Katta, Theo Coombes, Jenia Jitsev, and Aran Komatsuzaki. 2021. LAION-400M: open dataset of clip-filtered 400 million image-text pairs . <i>CoRR</i> , abs/2111.02114.	Hubert Siuzdak, Florian Grötschla, and Luca A. Lanzendörfer. 2024. SNAC: multi-scale neural audio codec . <i>CoRR</i> , abs/2410.14411.	2109 2110 2111
2053	Dustin Schwenk, Apoorv Khandelwal, Christopher Clark, Kenneth Marino, and Roozbeh Mottaghi. 2022. A-OKVQA: A benchmark for visual question answering using world knowledge . In <i>Computer Vision - ECCV 2022 - 17th European Conference, Tel Aviv, Israel, October 23-27, 2022, Proceedings, Part VIII</i> , volume 13668 of <i>Lecture Notes in Computer Science</i> , pages 146–162. Springer.	Shuran Song, Samuel P. Lichtenberg, and Jianxiong Xiao. 2015. SUN RGB-D: A RGB-D scene understanding benchmark suite . In <i>IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2015, Boston, MA, USA, June 7-12, 2015</i> , pages 567–576. IEEE Computer Society.	2112 2113 2114 2115 2116 2117
2061	Share. 2024. Sharegemin: Scaling up video caption data for multimodal large language models .	Khurram Soomro, Amir Roshan Zamir, and Mubarak Shah. 2012. UCF101: A dataset of 101 human actions classes from videos in the wild . <i>CoRR</i> , abs/1212.0402.	2118 2119 2120 2121
2063	Piyush Sharma, Nan Ding, Sebastian Goodman, and Radu Soricut. 2018. Conceptual captions: A cleaned, hypernymed, image alt-text dataset for automatic image captioning . In <i>Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics, ACL 2018, Melbourne, Australia, July 15-20, 2018, Volume 1: Long Papers</i> , pages 2556–2565. Association for Computational Linguistics.	Yixuan Su, Tian Lan, Huayang Li, Jialu Xu, Yan Wang, and Deng Cai. 2023. Pandagpt: One model to instruction-follow them all . <i>CoRR</i> , abs/2305.16355.	2122 2123 2124
2071	Yiming Shi, Xun Zhu, Ying Hu, Chenyi Guo, Miao Li, and Ji Wu. 2024. Med-2e3: A 2d-enhanced 3d medical multimodal large language model . <i>CoRR</i> , abs/2411.12783.	Guangzhi Sun, Wenyi Yu, Changli Tang, Xianzhao Chen, Tian Tan, Wei Li, Lu Lu, Zejun Ma, Yuxuan Wang, and Chao Zhang. 2024a. video-salmonn: Speech-enhanced audio-visual large language models . In <i>Forty-first International Conference on Machine Learning, ICML 2024, Vienna, Austria, July 21-27, 2024</i> . OpenReview.net.	2125 2126 2127 2128 2129 2130 2131

2132	Guangzhi Sun, Wenyi Yu, Changli Tang, Xianzhao	Zineng Tang, Ziyi Yang, Mahmoud Khademi, Yang Liu,	2190
2133	Chen, Tian Tan, Wei Li, Lu Lu, Zejun Ma, and Chao	Chenguang Zhu, and Mohit Bansal. 2024b. Codi-	2191
2134	Zhang. 2023a. Fine-grained audio-visual joint rep-	2: In-context, interleaved, and interactive any-to-any	2192
2135	representations for multimodal large language models.	generation . In <i>IEEE/CVF Conference on Computer</i>	2193
2136	<i>CoRR</i> , abs/2310.05863.	<i>Vision and Pattern Recognition, CVPR 2024, Seattle,</i>	2194
		WA, USA, June 16-22, 2024, pages 27415–27424.	2195
2137	Keqiang Sun, Junting Pan, Yuying Ge, Hao Li, Haodong	IEEE.	2196
2138	Duan, Xiaoshi Wu, Renrui Zhang, Aojun Zhou,		
2139	Zipeng Qin, Yi Wang, Jifeng Dai, Yu Qiao, Limin	Zineng Tang, Ziyi Yang, Chenguang Zhu, Michael Zeng,	2197
2140	Wang, and Hongsheng Li. 2023b. Journeydb: A	and Mohit Bansal. 2023. Any-to-any generation via	2198
2141	benchmark for generative image understanding . In	composable diffusion . In <i>Advances in Neural In-</i>	2199
2142	<i>Advances in Neural Information Processing Systems</i>	<i>formation Processing Systems 36: Annual Confer-</i>	2200
2143	<i>36: Annual Conference on Neural Information Pro-</i>	<i>ence on Neural Information Processing Systems 2023,</i>	2201
2144	<i>cessing Systems 2023, NeurIPS 2023, New Orleans,</i>	<i>NeurIPS 2023, New Orleans, LA, USA, December 10</i>	2202
2145	<i>LA, USA, December 10 - 16, 2023.</i>	<i>- 16, 2023.</i>	2203
2146	Licai Sun, Zheng Lian, Bin Liu, and Jianhua Tao. 2023c.	Yapeng Tian, Dingzeyu Li, and Chenliang Xu. 2020.	2204
2147	MAE-DFER: efficient masked autoencoder for self-	Unified multisensory perception: Weakly-supervised	2205
2148	supervised dynamic facial expression recognition . In	audio-visual video parsing . In <i>Computer Vision -</i>	2206
2149	<i>Proceedings of the 31st ACM International Confer-</i>	<i>ECCV 2020 - 16th European Conference, Glasgow,</i>	2207
2150	<i>ence on Multimedia, MM 2023, Ottawa, ON, Canada,</i>	UK, August 23-28, 2020, <i>Proceedings, Part III</i> , vol-	2208
2151	<i>29 October 2023- 3 November 2023</i> , pages 6110–	ume 12348 of <i>Lecture Notes in Computer Science</i> ,	2209
2152	6121. ACM.	pages 436–454. Springer.	2210
2153	Licai Sun, Zheng Lian, Bin Liu, and Jianhua Tao. 2023d.	Zhan Tong, Yibing Song, Jue Wang, and Limin	2211
2154	MAE-DFER: efficient masked autoencoder for self-	Wang. 2022. Videomae: Masked autoencoders are	2212
2155	supervised dynamic facial expression recognition . In	data-efficient learners for self-supervised video pre-	2213
2156	<i>Proceedings of the 31st ACM International Confer-</i>	training . In <i>Advances in Neural Information Pro-</i>	2214
2157	<i>ence on Multimedia, MM 2023, Ottawa, ON, Canada,</i>	<i>cessing Systems 35: Annual Conference on Neural</i>	2215
2158	<i>29 October 2023- 3 November 2023</i> , pages 6110–	<i>Information Processing Systems 2022, NeurIPS 2022,</i>	2216
2159	6121. ACM.	<i>New Orleans, LA, USA, November 28 - December 9,</i>	2217
		2022.	2218
2160	Luoyi Sun, Xuenan Xu, Mengyue Wu, and Weidi Xie.	Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier	2219
2161	2024b. Auto-acd: A large-scale dataset for audio-	Martinet, Marie-Anne Lachaux, Timothée Lacroix,	2220
2162	language representation learning . In <i>Proceedings</i>	Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal	2221
2163	<i>of the 32nd ACM International Conference on Mul-</i>	Azhar, Aurélien Rodriguez, Armand Joulin, Edouard	2222
2164	<i>timedia, MM 2024, Melbourne, VIC, Australia, 28</i>	Grave, and Guillaume Lample. 2023. Llama: Open	2223
2165	<i>October 2024 - 1 November 2024</i> , pages 5025–5034.	and efficient foundation language models . <i>CoRR</i> ,	2224
2166	ACM.	abs/2302.13971.	2225
2167	Quan Sun, Yuxin Fang, Ledell Wu, Xinlong Wang, and	Jack Urbanek, Florian Bordes, Pietro Astolfi, Mary	2226
2168	Yue Cao. 2023e. EVA-CLIP: improved training tech-	Williamson, Vasu Sharma, and Adriana Romero-	2227
2169	niques for CLIP at scale . <i>CoRR</i> , abs/2303.15389.	Soriano. 2024. A picture is worth more than 77 text	2228
		tokens: Evaluating clip-style models on dense cap-	2229
2170	Quan Sun, Qiyang Yu, Yufeng Cui, Fan Zhang,	tions . In <i>IEEE/CVF Conference on Computer Vision</i>	2230
2171	Xiaosong Zhang, Yuezhe Wang, Hongcheng Gao,	<i>and Pattern Recognition, CVPR 2024, Seattle, WA,</i>	2231
2172	Jingjing Liu, Tiejun Huang, and Xinlong Wang.	USA, June 16-22, 2024, pages 26690–26699. IEEE.	2232
2173	2024c. Emu: Generative pretraining in multimodal-		
2174	ity . In <i>The Twelfth International Conference on</i>	Christophe Veaux, Junichi Yamagishi, Kirsten MacDon-	2233
2175	<i>Learning Representations, ICLR 2024, Vienna, Aus-</i>	ald, et al. 2017. Cstr vctk corpus: English multi-	2234
2176	<i>tria, May 7-11, 2024</i> . OpenReview.net.	speaker corpus for cstr voice cloning toolkit. <i>CSTR</i> ,	2235
		6:15.	2236
2177	Andrew Szot, Bogdan Mazouze, Harsh Agrawal, R. De-		
2178	von Hjelm, Zsolt Kira, and Alexander Toshev. 2024a.	Andreas Veit, Tomas Matera, Lukás Neumann, Jiri	2237
2179	Grounding multimodal large language models in ac-	Matas, and Serge J. Belongie. 2016. Coco-text:	2238
2180	tions . <i>CoRR</i> , abs/2406.07904.	Dataset and benchmark for text detection and recog-	2239
		nition in natural images . <i>CoRR</i> , abs/1601.07140.	2240
2181	Andrew Szot, Bogdan Mazouze, Omar Attia, Aleksei		
2182	Timofeev, Harsh Agrawal, Devon Hjelm, Zhe Gan,	Jiaqi Wang, Hanqi Jiang, Yiheng Liu, Chong Ma,	2241
2183	Zsolt Kira, and Alexander Toshev. 2024b. From mul-	Xu Zhang, Yi Pan, Mengyuan Liu, Peiran Gu,	2242
2184	timodal llms to generalist embodied agents: Methods	Sichen Xia, Wenjun Li, Yutong Zhang, Zihao Wu,	2243
2185	and lessons . <i>Preprint</i> , arXiv:2412.08442.	Zhengliang Liu, Tianyang Zhong, Bao Ge, Tuo	2244
		Zhang, Ning Qiang, Xintao Hu, Xi Jiang, Xin Zhang,	2245
2186	Yunlong Tang, Daiki Shimada, Jing Bi, and Chenliang	Wei Zhang, Dinggang Shen, Tianming Liu, and Shu	2246
2187	Xu. 2024a. Avicuna: Audio-visual LLM with inter-		
2188	leaver and context-boundary alignment for temporal		
2189	referential dialogue . <i>CoRR</i> , abs/2403.16276.		

2247	Zhang. 2024a. A comprehensive review of multi-modal large language models: Performance and challenges across different tasks . <i>CoRR</i> , abs/2408.01319.	2305
2248		2306
2249		2307
2250	Jiuniu Wang, Hangjie Yuan, Dayou Chen, Yingya Zhang, Xiang Wang, and Shiwei Zhang. 2023a. Modelscape text-to-video technical report . <i>CoRR</i> , abs/2308.06571.	2308
2251		2309
2252		2310
2253		2311
2254	Kai Wang, Boris Babenko, and Serge J. Belongie. 2011. End-to-end scene text recognition . In <i>IEEE International Conference on Computer Vision, ICCV 2011, Barcelona, Spain, November 6-13, 2011</i> , pages 1457–1464. IEEE Computer Society.	2312
2255		2313
2256		2314
2257		2315
2258		2316
2259	Peng Wang, Shuai Bai, Sinan Tan, Shijie Wang, Zhihao Fan, Jinze Bai, Keqin Chen, Xuejing Liu, Jialin Wang, Wenbin Ge, Yang Fan, Kai Dang, Mengfei Du, Xuancheng Ren, Rui Men, Dayiheng Liu, Chang Zhou, Jingren Zhou, and Junyang Lin. 2024b. Qwen2-vl: Enhancing vision-language model’s perception of the world at any resolution . <i>CoRR</i> , abs/2409.12191.	2317
2260		2318
2261		2319
2262		2320
2263		2321
2264		2322
2265		2323
2266		2324
2267	Wenhai Wang, Jiangwei Xie, Chuanyang Hu, Haoming Zou, Jianan Fan, Wenwen Tong, Yang Wen, Silei Wu, Hanming Deng, Zhiqi Li, Hao Tian, Lewei Lu, Xizhou Zhu, Xiaogang Wang, Yu Qiao, and Jifeng Dai. 2023b. Drivemlm: Aligning multi-modal large language models with behavioral planning states for autonomous driving . <i>CoRR</i> , abs/2312.09245.	2325
2268		2326
2269		2327
2270		2328
2271		2329
2272		2330
2273		2331
2274	Xin Wang, Jiawei Wu, Junkun Chen, Lei Li, Yuanfang Wang, and William Yang Wang. 2019. Vatex: A large-scale, high-quality multilingual dataset for video-and-language research . In <i>2019 IEEE/CVF International Conference on Computer Vision, ICCV 2019, Seoul, Korea (South), October 27 - November 2, 2019</i> , pages 4580–4590. IEEE.	2332
2275		2333
2276		2334
2277		2335
2278		2336
2279		2337
2280		2338
2281	Xinyu Wang, Bohan Zhuang, and Qi Wu. 2024c. Modaverse: Efficiently transforming modalities with llms . In <i>IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2024, Seattle, WA, USA, June 16-22, 2024</i> , pages 26596–26606. IEEE.	2339
2282		2340
2283		2341
2284		2342
2285		2343
2286	Yi Wang, Yanan He, Yizhuo Li, Kunchang Li, Jiashuo Yu, Xin Ma, Xinhao Li, Guo Chen, Xinyuan Chen, Yaohui Wang, Ping Luo, Ziwei Liu, Yali Wang, Limin Wang, and Yu Qiao. 2024d. Internvid: A large-scale video-text dataset for multimodal understanding and generation . In <i>The Twelfth International Conference on Learning Representations, ICLR 2024, Vienna, Austria, May 7-11, 2024</i> . OpenReview.net.	2344
2287		2345
2288		2346
2289		2347
2290		2348
2291		2349
2292		2350
2293		2351
2294	Yi Wang, Kunchang Li, Xinhao Li, Jiashuo Yu, Yanan He, Guo Chen, Baoqi Pei, Rongkun Zheng, Zun Wang, Yansong Shi, Tianxiang Jiang, Songze Li, Jilan Xu, Hongjie Zhang, Yifei Huang, Yu Qiao, Yali Wang, and Limin Wang. 2024e. Internvideo2: Scaling foundation models for multimodal video understanding . In <i>Computer Vision - ECCV 2024 - 18th European Conference, Milan, Italy, September 29-October 4, 2024, Proceedings, Part LXXXV</i> , volume 15143 of <i>Lecture Notes in Computer Science</i> , pages 396–416. Springer.	2352
2295		2353
2296		2354
2297		2355
2298		2356
2299		2357
2300		2358
2301		2359
2302		2360
2303		2361
2304		
	Zekun Wang, King Zhu, Chunpu Xu, Wangchunshu Zhou, Jiaheng Liu, Yibo Zhang, Jiashuo Wang, Ning Shi, Siyu Li, Yizhi Li, Haoran Que, Zhaoxiang Zhang, Yuanxing Zhang, Ge Zhang, Ke Xu, Jie Fu, and Wenhao Huang. 2024f. Mio: A foundation model on multimodal tokens . <i>Preprint</i> , arXiv:2409.17692.	
	Julong Wei, Shanshuai Yuan, Pengfei Li, Qingda Hu, Zhongxue Gan, and Wenchao Ding. 2024. Occllama: An occupancy-language-action generative world model for autonomous driving . <i>CoRR</i> , abs/2409.03272.	
	Bo Wu, Shoubin Yu, Zhenfang Chen, Josh Tenenbaum, and Chuang Gan. 2021. STAR: A benchmark for situated reasoning in real-world videos . In <i>Proceedings of the Neural Information Processing Systems Track on Datasets and Benchmarks 1, NeurIPS Datasets and Benchmarks 2021, December 2021, virtual</i> .	
	Jiahong Wu, He Zheng, Bo Zhao, Yixin Li, Baoming Yan, Rui Liang, Wenjia Wang, Shipai Zhou, Guosen Lin, Yanwei Fu, Yizhou Wang, and Yonggang Wang. 2017. AI challenger : A large-scale dataset for going deeper in image understanding . <i>CoRR</i> , abs/1711.06475.	
	Jialian Wu, Jianfeng Wang, Zhengyuan Yang, Zhe Gan, Zicheng Liu, Junsong Yuan, and Lijuan Wang. 2024a. Grit: A generative region-to-text transformer for object understanding . In <i>Computer Vision - ECCV 2024 - 18th European Conference, Milan, Italy, September 29-October 4, 2024, Proceedings, Part LXXX</i> , volume 15138 of <i>Lecture Notes in Computer Science</i> , pages 207–224. Springer.	
	Jiayang Wu, Wensheng Gan, Zefeng Chen, Shicheng Wan, and Philip S. Yu. 2023. Multimodal large language models: A survey . In <i>IEEE International Conference on Big Data, BigData 2023, Sorrento, Italy, December 15-18, 2023</i> , pages 2247–2256. IEEE.	
	Shengqiong Wu, Hao Fei, Leigang Qu, Wei Ji, and Tat-Seng Chua. 2024b. Next-gpt: Any-to-any multimodal LLM . In <i>Forty-first International Conference on Machine Learning, ICML 2024, Vienna, Austria, July 21-27, 2024</i> . OpenReview.net.	
	Zhirong Wu, Shuran Song, Aditya Khosla, Fisher Yu, Linguang Zhang, Xiaoou Tang, and Jianxiong Xiao. 2015. 3d shapenets: A deep representation for volumetric shapes . In <i>IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2015, Boston, MA, USA, June 7-12, 2015</i> , pages 1912–1920. IEEE Computer Society.	
	Hanguang Xiao, Feizhong Zhou, Xingyue Liu, Tianqi Liu, Zhipeng Li, Xin Liu, and Xiaoxuan Huang. 2024. A comprehensive survey of large language models and multimodal large language models in medicine . <i>CoRR</i> , abs/2405.08603.	
	Junbin Xiao, Xindi Shang, Angela Yao, and Tat-Seng Chua. 2021. Next-qa: Next phase of question-answering to explaining temporal actions . In <i>IEEE</i>	

2362	<i>Conference on Computer Vision and Pattern Recognition, CVPR 2021, virtual, June 19-25, 2021, pages 9777–9786. Computer Vision Foundation / IEEE.</i>	2418
2363		2419
2364		2420
2365	Zhifei Xie and Changqiao Wu. 2024. Mini-omni2: Towards open-source gpt-4o with vision, speech and duplex capabilities. <i>CoRR</i> , abs/2410.11190.	2421
2366		2422
2367		2423
2368	Dejing Xu, Zhou Zhao, Jun Xiao, Fei Wu, Hanwang Zhang, Xiangnan He, and Yueting Zhuang. 2017. Video question answering via gradually refined attention over appearance and motion. In <i>Proceedings of the 2017 ACM on Multimedia Conference, MM 2017, Mountain View, CA, USA, October 23-27, 2017</i> , pages 1645–1653. ACM.	2424
2369		2425
2370		2426
2371		2427
2372		2428
2373		2429
2374		2430
2375	Haofei Xu, Jing Zhang, Jianfei Cai, Hamid Rezaatofghi, Fisher Yu, Dacheng Tao, and Andreas Geiger. 2023. Unifying flow, stereo and depth estimation. <i>IEEE Trans. Pattern Anal. Mach. Intell.</i> , 45(11):13941–13958.	2431
2376		2432
2377		2433
2378		2434
2379		2435
2380	Jingwei Xu, Chenyu Wang, Zibo Zhao, Wen Liu, Yi Ma, and Shenghua Gao. 2024a. CAD-MLLM: unifying multimodality-conditioned CAD generation with MLLM. <i>CoRR</i> , abs/2411.04954.	2436
2381		2437
2382		2438
2383		2439
2384	Jun Xu, Tao Mei, Ting Yao, and Yong Rui. 2016. MSR-VTT: A large video description dataset for bridging video and language. In <i>2016 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2016, Las Vegas, NV, USA, June 27-30, 2016</i> , pages 5288–5296. IEEE Computer Society.	2440
2385		2441
2386		2442
2387		2443
2388		2444
2389		2445
2390	Runsun Xu, Xiaolong Wang, Tai Wang, Yilun Chen, Jiangmiao Pang, and Dahua Lin. 2024b. Pointllm: Empowering large language models to understand point clouds. In <i>Computer Vision - ECCV 2024 - 18th European Conference, Milan, Italy, September 29-October 4, 2024, Proceedings, Part XXV</i> , volume 15083 of <i>Lecture Notes in Computer Science</i> , pages 131–147. Springer.	2446
2391		2447
2392		2448
2393		2449
2394		2450
2395		2451
2396		2452
2397		2453
2398	Jinlong Xue, Yayue Deng, Yingming Gao, and Ya Li. 2024a. Auffusion: Leveraging the power of diffusion and large language models for text-to-audio generation. <i>IEEE ACM Trans. Audio Speech Lang. Process.</i> , 32:4700–4712.	2454
2399		2455
2400		2456
2401		2457
2402		2458
2403	Le Xue, Ning Yu, Shu Zhang, Artemis Panagopoulou, Junnan Li, Roberto Martín-Martín, Jiajun Wu, Caiming Xiong, Ran Xu, Juan Carlos Niebles, and Silvio Savarese. 2024b. ULIP-2: towards scalable multimodal pre-training for 3d understanding. In <i>IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2024, Seattle, WA, USA, June 16-22, 2024</i> , pages 27081–27091. IEEE.	2459
2404		2460
2405		2461
2406		2462
2407		2463
2408		2464
2409		2465
2410		2466
2411	Honghui Yang, Tong He, Jiaheng Liu, Hua Chen, Boxi Wu, Binbin Lin, Xiaofei He, and Wanli Ouyang. 2023a. GD-MAE: generative decoder for MAE pre-training on lidar point clouds. In <i>IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2023, Vancouver, BC, Canada, June 17-24, 2023</i> , pages 9403–9414. IEEE.	2467
2412		2468
2413		2469
2414		2470
2415		2471
2416		2472
2417		2473
		2474
		2475
	Zhen Yang, Yingxue Zhang, Fandong Meng, and Jie Zhou. 2023b. TEAL: tokenize and embed ALL for multi-modal large language models. <i>CoRR</i> , abs/2311.04589.	
	Hanrong Ye, De-An Huang, Yao Lu, Zhiding Yu, Wei Ping, Andrew Tao, Jan Kautz, Song Han, Dan Xu, Pavlo Molchanov, and Hongxu Yin. 2024a. X-VILA: cross-modality alignment for large language model. <i>CoRR</i> , abs/2405.19335.	
	Qilang Ye, Zitong Yu, Rui Shao, Xinyu Xie, Philip Torr, and Xiaochun Cao. 2024b. CAT: enhancing multi-modal large language model to answer questions in dynamic audio-visual scenarios. In <i>Computer Vision - ECCV 2024 - 18th European Conference, Milan, Italy, September 29-October 4, 2024, Proceedings, Part X</i> , volume 15068 of <i>Lecture Notes in Computer Science</i> , pages 146–164. Springer.	
	Shukang Yin, Chaoyou Fu, Sirui Zhao, Ke Li, Xing Sun, Tong Xu, and Enhong Chen. 2023a. A survey on multimodal large language models. <i>CoRR</i> , abs/2306.13549.	
	Zhenfei Yin, Jiong Wang, Jianjian Cao, Zhelun Shi, Dingning Liu, Mukai Li, Xiaoshui Huang, Zhiyong Wang, Lu Sheng, Lei Bai, Jing Shao, and Wanli Ouyang. 2023b. LAMM: language-assisted multi-modal instruction-tuning dataset, framework, and benchmark. In <i>Advances in Neural Information Processing Systems 36: Annual Conference on Neural Information Processing Systems 2023, NeurIPS 2023, New Orleans, LA, USA, December 10 - 16, 2023</i> .	
	Alex Young, Bei Chen, Chao Li, Chengen Huang, Ge Zhang, Guanwei Zhang, Heng Li, Jiangcheng Zhu, Jianqun Chen, Jing Chang, Kaidong Yu, Peng Liu, Qiang Liu, Shawn Yue, Senbin Yang, Shiming Yang, Tao Yu, Wen Xie, Wenhao Huang, Xiaohui Hu, Xiaoyi Ren, Xinyao Niu, Pengcheng Nie, Yuchi Xu, Yudong Liu, Yue Wang, Yuxuan Cai, Zhenyu Gu, Zhiyuan Liu, and Zonghong Dai. 2024. Yi: Open foundation models by 01.ai. <i>CoRR</i> , abs/2403.04652.	
	Jiazuo Yu, Haomiao Xiong, Lu Zhang, Haiwen Diao, Yunzhi Zhuge, Lanqing Hong, Dong Wang, Huchuan Lu, You He, and Long Chen. 2024a. Llms can evolve continually on modality for x-modal reasoning. In <i>Advances in Neural Information Processing Systems 38: Annual Conference on Neural Information Processing Systems 2024, NeurIPS 2024, Vancouver, BC, Canada, December 10 - 15, 2024</i> .	
	Licheng Yu, Patrick Poirson, Shan Yang, Alexander C. Berg, and Tamara L. Berg. 2016. Modeling context in referring expressions. In <i>Computer Vision - ECCV 2016 - 14th European Conference, Amsterdam, The Netherlands, October 11-14, 2016, Proceedings, Part II</i> , volume 9906 of <i>Lecture Notes in Computer Science</i> , pages 69–85. Springer.	
	Lijun Yu, José Lezama, Nitesh Bharadwaj Gundavarapu, Luca Versari, Kihyuk Sohn, David Minnen, Yong Cheng, Agrim Gupta, Xiuye Gu, Alexander G. Hauptmann, Boqing Gong, Ming-Hsuan Yang, Irfan Essa,	

2476	David A. Ross, and Lu Jiang. 2024b. Language model beats diffusion - tokenizer is key to visual generation . In <i>The Twelfth International Conference on Learning Representations, ICLR 2024, Vienna, Austria, May 7-11, 2024</i> . OpenReview.net.	2534
2477		2535
2478		2536
2479		2537
2480		2538
2481	Shoubin Yu, Jaehong Yoon, and Mohit Bansal. 2024c. Crema: Generalizable and efficient video-language reasoning via multimodal modular fusion . <i>Preprint</i> , arXiv:2402.05889.	2539
2482		2540
2483		2541
2484		2542
2485		2543
2486	Weihao Yu, Zhengyuan Yang, Linjie Li, Jianfeng Wang, Kevin Lin, Zicheng Liu, Xinchao Wang, and Lijuan Wang. 2024d. Mm-vet: Evaluating large multimodal models for integrated capabilities . In <i>Forty-first International Conference on Machine Learning, ICML 2024, Vienna, Austria, July 21-27, 2024</i> . OpenReview.net.	2544
2487		2545
2488		2546
2489		2547
2490		2548
2491		2549
2492		
2493	Zhou Yu, Dejing Xu, Jun Yu, Ting Yu, Zhou Zhao, Yuet-ing Zhuang, and Dacheng Tao. 2019. Activitynet-qa: A dataset for understanding complex web videos via question answering . In <i>The Thirty-Third AAAI Conference on Artificial Intelligence, AAAI 2019, The Thirty-First Innovative Applications of Artificial Intelligence Conference, IAAI 2019, The Ninth AAAI Symposium on Educational Advances in Artificial Intelligence, EAAI 2019, Honolulu, Hawaii, USA, January 27 - February 1, 2019</i> , pages 9127–9134. AAAI Press.	2550
2494		2551
2495		2552
2496		2553
2497		2554
2498		2555
2499		2556
2500		2557
2501		2558
2502		2559
2503	Xiang Yue, Yuansheng Ni, Tianyu Zheng, Kai Zhang, Ruoqi Liu, Ge Zhang, Samuel Stevens, Dongfu Jiang, Weiming Ren, Yuxuan Sun, Cong Wei, Botao Yu, Ruibin Yuan, Renliang Sun, Ming Yin, Boyuan Zheng, Zhenzhu Yang, Yibo Liu, Wenhao Huang, Huan Sun, Yu Su, and Wenhui Chen. 2024. MMMUI: A massive multi-discipline multimodal understanding and reasoning benchmark for expert AGI . In <i>IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2024, Seattle, WA, USA, June 16-22, 2024</i> , pages 9556–9567. IEEE.	2560
2504		2561
2505		2562
2506		2563
2507		2564
2508		2565
2509		2566
2510		2567
2511		
2512		2568
2513		2569
2514	Neil Zeghidour, Alejandro Luebs, Ahmed Omran, Jan Skoglund, and Marco Tagliasacchi. 2022. Soundstream: An end-to-end neural audio codec . <i>IEEE ACM Trans. Audio Speech Lang. Process.</i> , 30:495–507.	2570
2515		2571
2516		2572
2517		2573
2518		2574
2519		2575
2520	Rowan Zellers, Jiasen Lu, Ximing Lu, Youngjae Yu, Yanpeng Zhao, Mohammadreza Salehi, Aditya Kusupati, Jack Hessel, Ali Farhadi, and Yejin Choi. 2022. MERLOT RESERVE: neural script knowledge through vision and language and sound . In <i>IEEE/CVF Conference on Computer Vision and Pattern Recognition, CVPR 2022, New Orleans, LA, USA, June 18-24, 2022</i> , pages 16354–16366. IEEE.	2576
2521		2577
2522		2578
2523		2579
2524		2580
2525		2581
2526		2582
2527	Aohan Zeng, Bin Xu, Bowen Wang, Chenhui Zhang, Da Yin, Diego Rojas, Guanyu Feng, Hanlin Zhao, Hanyu Lai, Hao Yu, Hongning Wang, Jiadai Sun, Jiajie Zhang, Jiale Cheng, Jiayi Gui, Jie Tang, Jing Zhang, Juanzi Li, Lei Zhao, Lindong Wu, Lucen Zhong, Mingdao Liu, Minlie Huang, Peng Zhang, Qinkai Zheng, Rui Lu, Shuaiqi Duan, Shudan Zhang,	2583
2528		2584
2529		2585
2530		2586
2531		2587
2532		2588
2533		2589
		2590
	Shulin Cao, Shuxun Yang, Weng Lam Tam, Wenyi Zhao, Xiao Liu, Xiao Xia, Xiaohan Zhang, Xiaotao Gu, Xin Lv, Xinghan Liu, Xinyi Liu, Xinyue Yang, Xixuan Song, Xunkai Zhang, Yifan An, Yifan Xu, Yilin Niu, Yuntao Yang, Yueyan Li, Yushi Bai, Yuxiao Dong, Zehan Qi, Zhaoyu Wang, Zhen Yang, Zhengxiao Du, Zhenyu Hou, and Zihan Wang. 2024. Chatglm: A family of large language models from GLM-130B to GLM-4 all tools . <i>CoRR</i> , abs/2406.12793.	
	Xiaohua Zhai, Basil Mustafa, Alexander Kolesnikov, and Lucas Beyer. 2023. Sigmoid loss for language image pre-training . In <i>IEEE/CVF International Conference on Computer Vision, ICCV 2023, Paris, France, October 1-6, 2023</i> , pages 11941–11952. IEEE.	
	Jun Zhan, Junqi Dai, Jiasheng Ye, Yunhua Zhou, Dong Zhang, Zhigeng Liu, Xin Zhang, Ruibin Yuan, Ge Zhang, Linyang Li, Hang Yan, Jie Fu, Tao Gui, Tianxiang Sun, Yu-Gang Jiang, and Xipeng Qiu. 2024. Anygpt: Unified multimodal LLM with discrete sequence modeling . In <i>Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers), ACL 2024, Bangkok, Thailand, August 11-16, 2024</i> , pages 9637–9662. Association for Computational Linguistics.	
	Binbin Zhang, Hang Lv, Pengcheng Guo, Qijie Shao, Chao Yang, Lei Xie, Xin Xu, Hui Bu, Xiaoyu Chen, Chenchen Zeng, Di Wu, and Zhendong Peng. 2022a. WENETSPEECH: A 10000+ hours multi-domain mandarin corpus for speech recognition . In <i>IEEE International Conference on Acoustics, Speech and Signal Processing, ICASSP 2022, Virtual and Singapore, 23-27 May 2022</i> , pages 6182–6186. IEEE.	
	Dong Zhang, Shimin Li, Xin Zhang, Jun Zhan, Pengyu Wang, Yaqian Zhou, and Xipeng Qiu. 2023a. Speechgpt: Empowering large language models with intrinsic cross-modal conversational abilities . In <i>Findings of the Association for Computational Linguistics: EMNLP 2023, Singapore, December 6-10, 2023</i> , pages 15757–15773. Association for Computational Linguistics.	
	Duzhen Zhang, Yahan Yu, Jiahua Dong, Chenxing Li, Dan Su, Chenhui Chu, and Dong Yu. 2024a. Mm-llms: Recent advances in multimodal large language models . In <i>Findings of the Association for Computational Linguistics, ACL 2024, Bangkok, Thailand and virtual meeting, August 11-16, 2024</i> , pages 12401–12430. Association for Computational Linguistics.	
	Hang Zhang, Xin Li, and Lidong Bing. 2023b. Video-llama: An instruction-tuned audio-visual language model for video understanding . In <i>Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing, EMNLP 2023 - System Demonstrations, Singapore, December 6-10, 2023</i> , pages 543–553. Association for Computational Linguistics.	

2591	Meishan Zhang, Hao Fei, Bin Wang, Shengqiong Wu,	text . In <i>Advances in Neural Information Processing</i>	2648
2592	Yixin Cao, Fei Li, and Min Zhang. 2024b. Recogniz-	<i>Systems 36: Annual Conference on Neural Informa-</i>	2649
2593	ing everything from all modalities at once: Grounded	<i>tion Processing Systems 2023, NeurIPS 2023, New</i>	2650
2594	multimodal universal information extraction . In <i>Find-</i>	<i>Orleans, LA, USA, December 10 - 16, 2023.</i>	2651
2595	<i>ings of the Association for Computational Linguistics,</i>		
2596	<i>ACL 2024, Bangkok, Thailand and virtual meeting,</i>	Wanrong Zhu, Jack Hessel, Anas Awadalla,	2652
2597	<i>August 11-16, 2024</i> , pages 14498–14511. Associa-	Samir Yitzhak Gadre, Jesse Dodge, Alex Fang,	2653
2598	tion for Computational Linguistics.	Youngjae Yu, Ludwig Schmidt, William Yang Wang,	2654
		and Yejin Choi. 2023b. Multimodal C4: an open,	2655
2599	Susan Zhang, Stephen Roller, Naman Goyal, Mikel	billion-scale corpus of images interleaved with	2656
2600	Artetxe, Moya Chen, Shuohui Chen, Christopher	text . In <i>Advances in Neural Information Processing</i>	2657
2601	Dewan, Mona T. Diab, Xian Li, Xi Victoria Lin,	<i>Systems 36: Annual Conference on Neural Informa-</i>	2658
2602	Todor Mihaylov, Myle Ott, Sam Shleifer, Kurt Shus-	<i>tion Processing Systems 2023, NeurIPS 2023, New</i>	2659
2603	ter, Daniel Simig, Punit Singh Koura, Anjali Srid-	<i>Orleans, LA, USA, December 10 - 16, 2023.</i>	2660
2604	har, Tianlu Wang, and Luke Zettlemoyer. 2022b.		
2605	OPT: open pre-trained transformer language mod-		
2606	els . <i>CoRR</i> , abs/2205.01068.		
2607	Xin Zhang, Dong Zhang, Shimin Li, Yaqian Zhou, and		
2608	Xipeng Qiu. 2023c. Speechtokenizer: Unified speech		
2609	tokenizer for speech large language models . <i>CoRR</i> ,		
2610	abs/2308.16692.		
2611	Yiyuan Zhang, Kaixiong Gong, Kaipeng Zhang, Hong-		
2612	sheng Li, Yu Qiao, Wanli Ouyang, and Xiangyu Yue.		
2613	2023d. Meta-transformer: A unified framework for		
2614	multimodal learning . <i>CoRR</i> , abs/2307.10802.		
2615	Yang Zhao, Zhijie Lin, Daquan Zhou, Zilong Huang,		
2616	Jiashi Feng, and Bingyi Kang. 2023a. Bubogpt: En-		
2617	abling visual grounding in multi-modal llms . <i>CoRR</i> ,		
2618	abs/2307.08581.		
2619	Zijia Zhao, Longteng Guo, Tongtian Yue, Sihan Chen,		
2620	Shuai Shao, Xinxin Zhu, Zehuan Yuan, and Jing		
2621	Liu. 2023b. Chatbridge: Bridging modalities with		
2622	large language model as a language catalyst . <i>CoRR</i> ,		
2623	abs/2305.16103.		
2624	Zhisheng Zhong, Chengyao Wang, Yuqi Liu, Senqiao		
2625	Yang, Longxiang Tang, Yuechen Zhang, Jingyao Li,		
2626	Tianyuan Qu, Yanwei Li, Yukang Chen, Shaozuo Yu,		
2627	Sitong Wu, Eric Lo, Shu Liu, and Jiaya Jia. 2024.		
2628	Lyra: An efficient and speech-centric framework for		
2629	omni-cognition . <i>Preprint</i> , arXiv:2412.09501.		
2630	Bin Zhu, Bin Lin, Munan Ning, Yang Yan, Jiaxi Cui,		
2631	Hongfa Wang, Yatian Pang, Wenhao Jiang, Junwu		
2632	Zhang, Zongwei Li, Caiwan Zhang, Zhifeng Li,		
2633	Wei Liu, and Li Yuan. 2024a. Languagebind: Ex-		
2634	tending video-language pretraining to n-modality by		
2635	language-based semantic alignment . In <i>The Twelfth</i>		
2636	<i>International Conference on Learning Representa-</i>		
2637	<i>tions, ICLR 2024, Vienna, Austria, May 7-11, 2024.</i>		
2638	OpenReview.net.		
2639	Hongyan Zhu, Shuai Qin, Min Su, Chengzhi Lin, Anjie		
2640	Li, and Junfeng Gao. 2024b. Harnessing large vision		
2641	and language models in agriculture: A review . <i>CoRR</i> ,		
2642	abs/2407.19679.		
2643	Wanrong Zhu, Jack Hessel, Anas Awadalla,		
2644	Samir Yitzhak Gadre, Jesse Dodge, Alex Fang,		
2645	Youngjae Yu, Ludwig Schmidt, William Yang Wang,		
2646	and Yejin Choi. 2023a. Multimodal C4: an open,		
2647	billion-scale corpus of images interleaved with		

A Related Survey

With the advent of MLLMs, there are several surveys detailing the current progress of MLLMs. Yin et al. (2023a); Wu et al. (2023); Caffagni et al. (2024) focus on the early Vision-MLLMs, while Çoban et al. (2024) and Ma et al. (2024) respectively summarize the Audio-MLLMs and 3D-MLLMs. Zhang et al. (2024a); Wang et al. (2024a) conduct an investigation into various Specific-MLLMs of different modalities. He et al. (2024); Chen et al. (2024b) discuss the expansion of MLLM’s generative capabilities. Some works discuss MLLMs in specific domains, such as medicine (Xiao et al., 2024), agriculture (Zhu et al., 2024b), and autonomous driving (Cui et al., 2024). Some works highlight some specific tasks such as safety (Fan et al., 2024), hallucination (Bai et al., 2024b), and acceleration (Zhu et al., 2024b). And Li and Lu (2024); Huang and Zhang (2024) focus on the evaluation of MLLM performance.

Distinct from the above-mentioned surveys, this paper focuses on MLLMs that align multiple non-linguistic modalities² with LLMs (Omni-MLLMs), enabling cross-modal understanding or cross-modal generation. As the first systematic survey on Omni-MLLMs, we hope our work will serve as an overview of this emerging direction, fostering future research in the field.

B Details about Omni-MLLMs architectures

Table 1 presents the details of the structure of mainstream Omni-MLLMs. We will list some of the pre-trained models used.

B.1 Modality Encoder

Visual Specific-Encoder ViT (Dosovitskiy et al., 2021), SigCLIP ViT (Zhai et al., 2023), CLIP ViT (Radford et al., 2021), EVA CLIP ViT (Sun et al., 2023e), InternViT (Chen et al., 2023g), DINOv2 ViT (Oquab et al., 2024), DFNCLIP ViT (Fang et al., 2024a), and OpenCLIP ConvNext (Liu et al., 2022b) encode images to obtain continuous features. TimeSformer (Bertasius et al., 2021), VideoMAE (Tong et al., 2022), MAE-DFer (Sun et al., 2023c), Omni-VL (Sun et al., 2023d), Video-Swin (Liu et al., 2022a), and Vivit (Arnab et al., 2021) encode videos to obtain continuous features.

²Since MLLMs capable of comprehending both video and imagery generally process video as multiple frames and employ a single vision encoder, we categorize them as Specific-MLLMs, i.e. Vision-MLLMs.

Audio Specific-Encoder AST (Gong et al., 2021), Beats (Chen et al., 2023d), Whisper (Radford et al., 2023), HuBert (Hsu et al., 2021), CLAP (Elizalde et al., 2023), Conformer (Gulati et al., 2020), MERT (Li et al., 2024d), and PANN (Kong et al., 2020) encode the audio modality to obtain continuous features.

3D Specific-Encoder ULIP2 (Xue et al., 2024b), GD-MAE (Yang et al., 2023a), PointEncoder (Xu et al., 2024b), FrozenCLIP (Huang et al., 2022), and M3D-CLIP (Bai et al., 2024a) encode the 3D modality to obtain continuous features.

Pre-align Uni-Encoder LanguageBind (Zhu et al., 2024a), ImageBind (Girdhar et al., 2023), Meta-Transformers (Zhang et al., 2023d), TVL (Fu et al., 2024c), and SSVTP (Kerr et al., 2023) encode multiple non-linguistic modalities into a unified feature space and obtain continuous features. TVL, LanguageBind, ImageBind, and SSVTP construct modality-specific encoders for different modalities and achieve multi-modalities alignment through indirect alignment. Meta-Transformers design distinct modality-specific patch embeddings and use a shared encoder to encode multiple modalities.

Other Specific-Encoder IMU2CLIP (Moon et al., 2022) encodes the IMU modality to obtain continuous features. Individual modality-specific encoders from LanguageBind or ImageBind are often used independently as specific encoders.

B.2 Modality Tokenizer

Visual Tokenizer VQ-GAN (Esser et al., 2021), DALL-E (Ramesh et al., 2021), BEiT-V2 (Peng et al., 2022), MAGVIT-v2 (Yu et al., 2024b), and SEED (Ge et al., 2023) encode the visual modality into discrete visual tokens, which can be decoded back into the original image using the de-tokenizer.

Audio Tokenizer Jukebox (Dhariwal et al., 2020), SoundStream (Zeghidour et al., 2022), SpeechTokenizer (Zhang et al., 2023c), Encoder (Défossez et al., 2023), and S2U (Chen et al., 2024c) encode the audio modality into discrete audio tokens, which can be decoded back into the audio using the corresponding de-tokenizer.

Other Tokenizer Scene Tokenizer (Wei et al., 2024) encodes the 3D modality into discrete 3D tokens. LEO (Huang et al., 2024), Ground-Action (Szot et al., 2024a), OccLLaMA (Wei et al.,

2024), and GMA (Szot et al., 2024b) perform discrete encoding of the action modality to obtain corresponding action tokens, which can be decoded back into the original action using the corresponding de-tokenizer. M3GPT (Luo et al., 2024), Gesticulator (Pang et al., 2024), and SOLAMI (Jiang et al., 2024b) perform discrete encoding of the motion modality to obtain corresponding motion tokens, which can be decoded back into the original motion using the corresponding de-tokenizer.

B.3 Modality Generation Model

For image generation, Stable Diffusion (Rombach et al., 2022) and Instruct-Pix2Pix (Brooks et al., 2023) are used. Video generation models include Zeroscope (Cerspense, 2023), VideoFusion (Luo et al., 2023c), VideoCrafter (Chen et al., 2023b), and ModelScope (Wang et al., 2023a). For audio generation, models such as AudioLDM (Liu et al., 2023b), SNAC (Siuzdak et al., 2024), LLaMA-Omni’s audio decoder (Fang et al., 2024b), MusicGen (Copet et al., 2023), and TiCodec (Ren et al., 2024) are utilized. Meanwhile, StyleTTS (Li et al., 2022d) and GPT-SoVITS (RVC-Boss) are employed for speech generation.

B.4 LLM Backbone

Commonly used LLMs include the T5 series (Raffel et al., 2020), LLaMA series (Touvron et al., 2023), Qwen series (Bai et al., 2023), Internlm series (Cai et al., 2024), Chatglm series (Zeng et al., 2024), OPT series (Zhang et al., 2022b), Mixtral series (Jiang et al., 2024a), Mistral series (Jiang et al., 2023), Phi series (Gunasekar et al., 2023), and Yi series (Young et al., 2024).

C Details of Training and evaluation

C.1 Details of Training Data

The statistical results of some commonly used alignment datasets and the instruction data of mainstream Omni-MLLMs are shown in Table 2 and Table 3. There is still a lack of alignment data for data-scarcity modalities and cross-modal instruction data.

C.2 Details of Benchmark

The statistical data of some commonly used benchmarks are shown in Table 4. Existing benchmarks still require improvements in terms of the number of modalities and the forms of modality interaction.

C.3 Other Train Recipes

In addition to the general training paradigms mentioned in Section 3, some other useful training recipes are also used. (1) **Prior knowledge from Specific-MLLMs**: Since Specific-MLLMs have already achieved effective alignment in single-modal scenarios, some Omni-MLLMs directly leverage their well-trained projectors to reduce the training overhead during the alignment phase. For example, InstructBLIP (Panagopoulou et al., 2024) and X-LLM (Chen et al., 2023a) use the Q-former trained by BLIP2 to align the visual modality, while NaviveMC and DAMC (Chen et al., 2024a) further leverage projectors from multiple models to handle alignment for visual, audio, and 3D modalities separately; (2) **Additional human preference training**: Szot et al. (2024b) and Ye et al. (2024b) adopt HF training methods like PPO and ADPO to better align with human preferences; (3) **Modalities Blending**: During progressive alignment pre-training or multi-step instruction fine-tuning, some works (Han et al., 2024a; Li et al., 2024c; Chen et al., 2024c) mix previously trained modality data with the current new modality data for training to prevent catastrophic forgetting.

C.4 Performance of Omni-MLLMs

We statistic the performance of various mainstream Omni-MLLMs in uni-modal understanding, cross-modal understanding, and cross-modal, as shown in Table 5. We also show the performance of several Specific-MLLMs (Lin et al., 2024; Li et al., 2024a; Chu et al., 2023; Xu et al., 2024b; Sun et al., 2024c; Jin et al., 2024) on selected tasks for comparison. The results are mainly from corresponding papers (some results are used as baselines in other papers). It is worth noting that due to differences in the size and performance of the pre-trained models, Omni-MLLMs with the same backbone LLM may still not be fairly comparable. Therefore, this table only provides a rough trend of performance.

It can be seen from the table that most Omni-MLLMs still exhibit a significant performance gap in uni-modal understanding tasks compared to Specific-MLLMs. Meanwhile, in uni-modal generation tasks, models like AnyGPT and CoDi-2 have achieved performance close to or even surpassing Specific-MLLMs. Additionally, Omni-MLLMs are capable of performing cross-modal tasks that Specific-MLLMs cannot handle.

Model	Capabilities	Modalities	Multi-Modalities Encoding Method	Encoding Model	Method	Multi-Modalities Alignment Projector	Vocabulary	Multi-Modalities Interaction LLM	Modalities	Multi-Modalities Generation Method	Generation model
eP-LLM	Cross-modal Understanding	Visual/Audio	Continuous Encoding	ViT/TimeSformer/AST	multi-branch	Linear	–	injection	–	–	–
VALOR	Cross-modal Understanding	Visual/Audio	Continuous Encoding	ViT/VideoSwiFi/AST	multi-branch	MLP	–	injection	–	–	–
X-LLM	Cross-modal Understanding	Visual/Audio	Continuous Encoding	ViT/Conformer	multi-branch	Q-former+Linear	–	concatenate	–	–	–
ChatBridge	Cross-modal Understanding	Visual/Audio	Continuous Encoding	EVAL-CLIP ViT/Beats	multi-branch	Preceiver	–	concatenate	–	–	–
PandaGPT	Cross-modal Understanding	Visual/Audio/3D	Continuous Encoding	ImageBind	uni-branch	Linear	–	concatenate	–	–	–
VideoLLaMA	Cross-modal Understanding	Visual/Audio	Continuous Encoding	EVA CLIP ViT/ ImageBind-Audio	multi-branch	Q-former+Linear	–	concatenate	–	–	–
LAMM	Cross-modal Understanding	Visual/Audio	Continuous Encoding	CLIP ViT/ ViT	multi-branch	MLP	–	concatenate	–	–	–
Macra-LLM	Cross-modal Understanding	Visual/Audio	Continuous Encoding	ViT/Whisper	multi-branch	Cross-Attention	–	concatenate	–	–	–
BuboGPT	Cross-modal Understanding	Visual/Audio	Continuous Encoding	CLIP ViT/ImageBind-Audio	multi-branch	Q-former+Linear	–	concatenate	–	–	–
Teal	&& Generation	Visual/Audio	Discrete Encoding	VQ-GAN/ Whisper-K-means	embedding	–	Added Vocabulary	concatenate	Image	Modality-Token-based	VQGAN de-tokenizer
ImageBind-LLM	Cross-modal Understanding	Visual/Audio/3D	Continuous Encoding	ImageBind	uni-branch	MLP	–	injection	–	–	–
Nexi-GPT	Cross-modal Understanding	Visual/Audio	Continuous Encoding	ImageBind	multi-branch	Linear	–	concatenate	Image/Videos/Audio	Representation-based	StableDiffusion1.5/ Zeroscope/AudioLD
Ary-MAL	Cross-modal Understanding	Visual/Audio/IMU	Continuous Encoding	CLIP ViT/CLAP/ IMUCLIP	multi-branch	Preceiver	–	injection	–	–	–
FAVOR	Cross-modal Understanding	Visual/Audio	Continuous Encoding	EVA CLIP ViT/Whisper	multi-branch	Q-former+Linear	–	concatenate	–	–	–
Octavies	Cross-modal Understanding	Visual/3D	Continuous Encoding	CLIP ViT/Object-Aw-Scene	multi-branch	MLP	–	concatenate	–	–	–
LEO	&& Generation	Visual/3D/Action	Hybrid Encoding	OpenCLIP Convnext/PointNet++/ LEO's Action tokenizer	multi-branch	MLP/ Spatial Transformer	Overwrite Vocabulary	concatenate	Action	Modality-Token-based	LEO's Action de-tokenizer
CoDo-2	Cross-modal Understanding	Visual/Audio	Continuous Encoding	ImageBind	uni-branch	MLP	–	concatenate	–	–	–
X-InstructBLIP	Cross-modal Understanding	Visual/Audio/3D	Continuous Encoding	EVA CLIP ViT/Beats/ ULIP2	multi-branch	Q-former+Linear	–	concatenate	Image/Videos/Audio	Representation-based	StableDiffusion2.1/ Zeroscope/AudioLD2
One-LLM	Cross-modal Understanding	Visual/Audio/3D/IMU/MRIR	Continuous Encoding	Meta-transformer	uni-branch	Q-former+Linear	–	concatenate	–	–	–
AV-LLM	Cross-modal Understanding	Normal Map	Continuous Encoding	CLIP ViT/CLAP	multi-branch	Linear	–	concatenate	–	–	–
DriveMLM	Cross-modal Understanding	Visual/3D	Continuous Encoding	EVA CLIP ViT/GD-MAE	multi-branch	Q-former+Linear	–	concatenate	–	–	–
Osmi-3D	Cross-modal Understanding	Visual/3D	Continuous Encoding	CLIP ViT/PointNet++	multi-branch	MLP	–	concatenate	–	–	–
MotaVerse	&& Generation	Visual/Audio	Continuous Encoding	ImageBind	multi-branch	Linear	–	concatenate	Image/Videos/Audio	Text-based	StableDiffusion /AudioLD/VideoFusion
MultiPLY	Cross-modal Understanding	Visual/Audio/3D /thermal/touch	Continuous Encoding	CLIP ViT/CLAP /ConceptGraph	multi-branch	MLP/Linear	–	concatenate	–	–	–
CREMA	Cross-modal Understanding	Visual/Audio/3D	Continuous Encoding	EVA CLIP ViT+Linear	uni-branch	Q-former+Linear	–	concatenate	–	–	–
GroundingGPT	Cross-modal Understanding	/thermal/focus/optical	Continuous Encoding	Beats+Linear/ConceptFusion+Linear	multi-branch	Q-former+Linear/MLP	–	concatenate	–	–	–
DAMC	Cross-modal Understanding	Visual/Audio	Continuous Encoding	CLIP ViT/ImageBind-Audio	multi-branch	Q-former+Linear/MLP	–	concatenate	–	–	–
Ary-GPT	Cross-modal Understanding	Visual/Audio	Discrete Encoding	SEED tokenizer/ Speech tokenizer/Encodes	embedding	–	Extend Vocabulary	concatenate	Image/Speech/Music	–	SEED de-tokenizer/Speech de-tokenizer/Encodes de-tokenizer
TVL-LLaMA	Cross-modal Understanding	Visual/Touch	Continuous Encoding	TVL Encoders	uni-branch	MLP	–	injection	–	–	–
SSVTP-LLaMA	Cross-modal Understanding	Visual/Touch	Continuous Encoding	SSVTP Encoders	uni-branch	MLP	–	injection	–	–	–
CAT	Cross-modal Understanding	Visual/Audio	Continuous Encoding	ImageBind	multi-branch	Linear	–	concatenate	–	–	–
AVicuna	Cross-modal Understanding	Visual/Audio	Continuous Encoding	CLIP ViT/CLAP	multi-branch	MLP	–	concatenate	–	–	–
WordGPT	Cross-modal Understanding	Visual/Audio	Continuous Encoding	LanguageBind	uni-branch	Linear	–	concatenate	–	–	–
QaP	Cross-modal Understanding	Visual/Audio	Continuous Encoding	CLIP ViT/CLAP	multi-branch	dot attention+Linear	–	injection	Image/Videos/Audio	Representation-based	–
Uni-Moe	Cross-modal Understanding	Visual/Audio	Continuous Encoding	CLIP ViT/Beats	multi-branch	Q-former+Linear/MLP	–	concatenate	–	–	–
MGPT	Cross-modal Understanding	Audio/Motion	Discrete Encoding	Jakobus tokenizer /MGPT's Motion tokenizer	embedding	Extend Vocabulary	–	concatenate	Music/Motion	Modality-Token-based	Jakobus de-tokenizer /MGPT's Motion de-tokenizer
X-VILA	Cross-modal Understanding	Visual/Audio	Continuous Encoding	ImageBind	multi-branch	MLP	–	concatenate	Image/Videos/Audio	Representation-based	Stable Diffusion /AudioLD/VideoCrafter
REAMO	Cross-modal Understanding	Visual/Audio	Continuous Encoding	ImageBind	multi-branch	Linear	–	concatenate	–	–	–
VideoLLaMA2	Cross-modal Understanding	Visual/Audio	Continuous Encoding	CLIP ViT/Beats	multi-branch	MLP	–	concatenate	–	–	–
Ground-Action	Cross-modal Understanding	Visual/Action	Hybrid Encoding	CLIP ViT /Ground-Action action tokenizer	multi-branch	Preceiver	Overwrite Vocabulary	concatenate	Action	Modality-Token-based	Ground-Action action de-tokenizer
Emotee-LLaMA	Cross-modal Understanding	Visual/Audio	Continuous Encoding	CLIP ViT	multi-branch	MLP	–	concatenate	–	–	–
EmpathyEar	Cross-modal Understanding	Visual/Audio	Continuous Encoding	/Hugging Chinese	uni-branch	Linear	–	concatenate	Video/Audio	Text-based	StyleTTS2/EAT
video-SALMONN	&& Generation	Visual/Audio	Continuous Encoding	ViT/Beats	multi-branch	Q-former+Linear	–	concatenate	–	–	–
Meerkat	Cross-modal Understanding	Visual/Audio	Continuous Encoding	CLIP ViT/CLAP	multi-branch	MLP	–	concatenate	–	–	–
InterOmni	Cross-modal Understanding	Visual/Audio	Continuous Encoding	Inter ViT/Whisper	multi-branch	Linear	–	concatenate	–	–	–
SynsLM	Cross-modal Understanding	Visual/Audio	Hybrid Encoding	SigCLIP ViT /XLSR-K-means	multi-branch	MLP	Extend Vocabulary	concatenate	–	–	–
UniteMLLM	Cross-modal Understanding	Visual/Audio	Continuous Encoding	CLIP ViT /ImageBind-Audio	multi-branch	Q-former+Linear	–	concatenate	–	–	–
VITA	Cross-modal Understanding	Visual/Audio	Continuous Encoding	CLIP ViT	multi-branch	MLP	–	concatenate	Image/Videos/Audio	Text-based	Instruct-pix2pix/ Affusion/ModelScope
OcLLaMA	Cross-modal Understanding	3D/Action	Discrete Encoding	/VITA's Audio Encoder	embedding	–	Extend Vocabulary	concatenate	Speech	Text-based	GPT-soVITS
LLaMA-AVSR	Cross-modal Understanding	Visual/Audio	Continuous Encoding	/OcLLaMA's 3D tokenizer	multi-branch	MLP	–	concatenate	Action/3D	Modality-Token-based	OcLLaMA's 3D de-tokenizer
MIO	Cross-modal Understanding	Visual/Audio	Discrete Encoding	AV-Hallnet /Whisper	multi-branch	MLP	–	concatenate	–	–	–
EMOVA	Cross-modal Understanding	Visual/Audio	Discrete Encoding	SEED tokenizer /Speech tokenizer	embedding	–	Extend Vocabulary	concatenate	Image/Speech	Modality-Token-based	SEED de-tokenizer /Speech de-tokenizer
LLM Gesticulator	Cross-modal Understanding	Visual/Audio	Hybrid Encoding	/EMOVA's S2U tokenizer	multi-branch	C-Abstractor	Extend Vocabulary	concatenate	Speech	Modality-Token-based	EMOVA's S2U de-tokenizer
Baihuai-Omni	Cross-modal Understanding	Audio/Motion	Discrete Encoding	MotionRVQ tokenizer /Encodes tokenizer	embedding	–	Extend Vocabulary	concatenate	Audio/Motion	Modality-Token-based	MotionRVQ de-tokenizer /Encodes de-tokenizer
EGMI	Cross-modal Understanding	Audio/Motion	Continuous Encoding	SigCLIP ViT /Whisper	multi-branch	CNN+MLP/Conv-GMLP	–	concatenate	–	–	–
Dolphin	Cross-modal Understanding	Audio/Motion	Continuous Encoding	ImageBind	uni-branch	Linear	–	concatenate	Image/Audio	Representation-based	StableDiffusion /AudioLD
Mini-Omni2	Cross-modal Understanding	Visual/Audio	Continuous Encoding	CLIP ViT/ImageBind-Audio	multi-branch	MLP	–	concatenate	–	–	–
OMCAT	Cross-modal Understanding	Visual/Audio	Continuous Encoding	CLIP ViT/ImageBind-Audio	multi-branch	MLP	–	concatenate	–	–	–
PatWeave	Cross-modal Understanding	Visual/Whisper	Continuous Encoding	EVA CLIP ViT/Beats	uni-branch	Q-former+transformer	–	concatenate	Audio	Modality-Token-based	SNAC de-tokenizer
CAD-MLLM	Cross-modal Understanding	Visual/3D	Continuous Encoding	/ULIP2 DINO v2 /Michelangelo	multi-branch	Preceiver+Linear/Linear	–	concatenate	–	–	–
EAGLE	Cross-modal Understanding	Visual/Audio	Continuous Encoding	ImageBind	multi-branch	MLP	–	concatenate	–	–	–
Spider	Cross-modal Understanding	Visual/Audio	Continuous Encoding	ImageBind	multi-branch	MLP	–	concatenate	Image/Videos/Audio	Text-based	StableDiffusion /AudioLD/VideoCrafter
Med-2e3	Cross-modal Understanding	Visual/3D	Continuous Encoding	SigCLIP ViT/3D-CLIP	multi-branch	Q-former+Linear/MLP	–	concatenate	Image/Videos/Audio	Text-based	StableDiffusion /AudioLD/VideoCrafter
LongVALE-LLM	Cross-modal Understanding	Visual/Audio	Continuous Encoding	CLIP ViT/Beats/Whisper	multi-branch	MLP	–	concatenate	–	–	–
SOLAMI	Cross-modal Understanding	Audio/Motion	Discrete Encoding	SpeechTokenizer	embedding	–	Extend Vocabulary	concatenate	–	–	–
MaMo-LLaMA	Cross-modal Understanding	Visual/Audio	Continuous Encoding	/SOLAMI's MotionTokenizer	embedding	–	Extend Vocabulary	concatenate	Speech/Motion	Modality-Token-based	Speech de-tokenizer /SOLAMI's Motion de-tokenizer
GMA	Cross-modal Understanding	Visual/Action	Hybrid Encoding	ViT/ViT/EMERT	multi-branch	Conv+MLP/Conv+Ran+MLP	–	injection	–	–	–
Lyra	Cross-modal Understanding	Visual/Audio	Continuous Encoding	SigCLIP ViT /GMA's Action tokenizer	multi-branch	MLP	Overwrite Vocabulary	concatenate	Image/Audio	Representation-based	MusicGen
VITA-1.5	Cross-modal Understanding	Visual/Audio	Continuous Encoding	DFNCLIP ViT/Whisper	multi-branch	MLP	–	concatenate	Action	Modality-Token-based	GMA's Action de-tokenizer
Empathetic-LLM	Cross-modal Understanding	Visual/Audio	Continuous Encoding	InterViT/VITA's Audio Encoder	multi-branch	MLP/CNN+MLP	–	concatenate	Audio	Representation-based	LLaMA-Omni's audio decoder
	Cross-modal Understanding	Visual/Audio	Continuous Encoding	CLIP ViT/Whisper	multi-branch	Q-former+Linear	–	concatenate	Speech	Representation-based	TiCodec decoder

Table 1: The architectures of mainstream OmniMLLMs. The architectures of 70 Omni-MLLMs are displayed by encoding, alignment, interaction, and generation.

Name	Type	Modality	#Sample
MSCOCO (Lin et al., 2014)	X-Text	Image,Text	620K
Visual Genome (Krishna et al., 2017b)	X-Text	Image,Text	4.5M
Flickr30k (Plummer et al., 2015)	X-Text	Image,Text	158K
SBU (Ordonez et al., 2011)	X-Text	Image,Text	1M
DCI (Urbanek et al., 2024)	X-Text	Image,Text	7.8K
BLIP-Capfilt (Li et al., 2022c)	X-Text	Image,Text	129M
AI Challenger captions (Wu et al., 2017)	X-Text	Image,Text	1.5M
Wukong Captions (Gu et al., 2022)	X-Text	Image,Text	101M
CC12M (Changpinyo et al., 2021)	X-Text	Image,Text	12.4M
CC3M (Sharma et al., 2018)	X-Text	Image,Text	3.3M
LAION-5B (Schuhmann et al., 2022)	X-Text	Image,Text	5.9B
Redcaps (Desai et al., 2021)	X-Text	Image,Text	12M
LAION-COCO (Schuhmann et al., 2022b(b))	X-Text	Image,Text	600M
LAION-CAT (Radenovic et al., 2023)	X-Text	Image,Text	440M
LAION-AESTHETICS (Schuhmann et al., 2022b(a))	X-Text	Image,Text	120M
ShareGPT4V (Chen et al., 2024e)	X-Text	Image,Text	1.2M
LAION-115M (Schuhmann et al., 2021)	X-Text	Image,Text	115M
Journeydb (Sun et al., 2023b)	X-Text	Image,Text	4.4M
Multimodal c4 (Zhu et al., 2023b)	X-Text-X	Image,Text	43.3M
OBELICS (Laurençon et al., 2023)	X-Text-X	Image,Text	141M
Panda-70M (Chen et al., 2024f)	X-Text	Video,Text	70M
Webvid2M (Bain et al., 2021)	X-Text	Video,Text	2M
Valley-Pretrain-703k (Luo et al., 2023a)	X-Text	Video,Text	703K
Webvid10M (Bain et al., 2021)	X-Text	Video,Text	10M
YT-Temporal (Zellers et al., 2022)	X-Text	Video,Text	180M
ActivityNet Captions (Krishna et al., 2017a)	X-Text	Video,Text	100K
InterVid (Wang et al., 2024d)	X-Text	Video,Text	10M
MSRVTT (Xu et al., 2016)	X-Text	Video,Text	200K
ShareGemini (Share, 2024)	X-Text	Video,Text	530K
AudioSet (Gemmeke et al., 2017)	X-Text	Audio,Text	2.1M
Clotho (Drossos et al., 2020)	X-Text	Audio,Text	5k
Auto-ACD (Sun et al., 2024b)	X-Text	Audio,Text	1.5M
AudioCap (Kim et al., 2019)	X-Text	Audio,Text	46k
WavCaps (Mei et al., 2024)	X-Text	Audio,Text	403K
AISHELL-1 (Bu et al., 2017)	X-Text	Audio,Text	128K
AISHELL-2 (Du et al., 2018)	X-Text	Audio,Text	1M
Gigaspeech (Chen et al., 2021)	X-Text	Speech,Text	–
Common Voice (Ardila et al., 2020)	X-Text	Speech,Text	–
MLS (Pratap et al., 2020)	X-Text	Speech,Text	–
Music caption (Zhan et al., 2024)	X-Text	Music,Text	100M
Cap3D (Luo et al., 2023b)	X-Text	3D,Text	1M
Objaverse (Deitke et al., 2023)	X-Text	3D,Text	800K
ScanRefer (Chen et al., 2020a)	X-Text	3D,Text	51.5K
Normal Caption (Han et al., 2024a)	X-Text	Normal,Text	0.5M
Depth Caption (Han et al., 2024a)	X-Text	Depth,Text	0.5M
NSD (Allen et al., 2022)	X-Text	fMRI,Text	9K
Ego4d (Grauman et al., 2022)	X-Text	Video,IMU,Text	528k
PU-VALOR (Tang et al., 2024a)	X-Y-Text	Video,Audio,Text	114K
VALOR (Chen et al., 2023e)	X-Y-Text	Video,Audio,Text	16k
VAST (Chen et al., 2023f)	X-Y-Text	Video,Audio,Text	414k
VIDAL (Zhu et al., 2024a)	X-Y-Text	Video, Thermal, Depth, Audio	10M
TVL (Fu et al., 2024c)	X-Y-Text	Image,Touch,Text	44K
M3D-Cap (Bai et al., 2024a)	X-Y-Text	Image,3D,Text	115K

Table 2: **The statistics for alignment datasets in Omni-MLLMs**, including single non-linguistic modality text pairing data (X-Text), multiple non-linguistic modalities text pairing data (X-Text-Y), and single non-linguistic modality text interleaved data (X-Text-X).

Name	Source	Task	Modality	Construction Method	#Sample
XLLM's SFT (Chen et al., 2023a)	MiniGPT-4, AISHELL-2, VSDial-CN, ActivityNet Caps	Uni-Modal Understanding,Cross-Modal Understanding	Image,Video,Audio,Text	Template Instructionalization, T2X generation	10k
ChatBridge's SFT (Zhao et al., 2023b)	MSRVTT, AudioCaps, VQAv2, VG-QA...	Uni-Modal Understanding,Cross-Modal Understanding	Image,Video,Audio,Text	Template Instructionalization, GPT generation	4.4M+209k
Macaw-LLM (Lyu et al., 2023)	MSCOCO, Charades, AVSD, VG-QA...	Uni-Modal Understanding,Cross-Modal Understanding	Image,Video,Audio,Text	GPT generation	69K+50K
BuboGPT's SFT	LLaVA, Clotho, VGGSS	Uni-Modal Understanding,Cross-Modal Understanding	Image,Audio,Text	Template Instructionalization, GPT generation	196K
NextGPT's SFT (Wu et al., 2024b)	WebVid, CC3M, AudioCap, Youtube...	Uni-Modal Understanding,Cross-Modal Understanding, Uni-Modal Generation,Cross-Modal Generation	Image,Video,Audio,Text	Template Instructionalization, GPT generation+retrieval, T2X generation	20K
AnyMal's SFT (Moon et al., 2024)	–	Uni-Modal Understanding	Image,Video,Audio,Text	Manual Annotation, GPT generation	210K
FAVOR's SFT (Sun et al., 2023a)	LLaVA, MSCOCO, Ego4D, LibriSpeech...	Uni-Modal Understanding,Cross-Modal Understanding	Image,Video,Audio,Text	Template Instructionalization, GPT generation	–
LEO's SFT (Huang et al., 2024)	ScanQA, SQA3D, 3RScan, CLIPPort...	Uni-Modal Understanding,Cross-Modal Understanding	Image,3D,Text,Action	Template Instructionalization, GPT generation	220k
CoDi-2's SFT (Tang et al., 2024b)	MIMIC-IT, LAION-400M, AudioSet, Webvid...	Uni-Modal Understanding,Uni-Modal Generation	Image,Audio,Text	Template Instructionalization	–
X-InstructBLIP's SFT (Panagopoulou et al., 2024)	MSCOCO, Clotho, MSVD, Cap3D...	Uni-Modal Understanding	Image,Video,Audio,3D,Text	Template Instructionalization, GPT generation	1.6M
OneLLM's SFT (Han et al., 2024a)	LLaVA-150K, Clotho, Ego4D, NSD...	Uni-Modal Understanding	Image,Video,Audio,3D,ImU, Depth,fMRI,Normal,Text	Template Instructionalization, T2X generation	2M
AVLLM's SFT (Shu et al., 2023)	ACAV100M, VGGSound, WebVid2M, WavCaps...	Uni-Modal Understanding,Cross-Modal Understanding	Video,Audio,Text	GPT generation	1.4M
Uni-IO2's SFT (Lu et al., 2024a)	CC3M, AudioSet, Webvid3m, Omni3D...	Uni-Modal Understanding,Cross-Modal Understanding, Uni-Modal Generation,Cross-Modal Generation	Image,Video,Audio,Text	Template Instructionalization, GPT generation	775m
ModaVerse's SFT (Wang et al., 2024c)	–	Uni-Modal Understanding,Cross-Modal Understanding, Uni-Modal Generation	Image,Video,Audio,Text	GPT generation	2M
REAMO's SFT (Zhang et al., 2024b)	–	Uni-Modal Understanding,Cross-Modal Understanding	Image,Video,Audio,Text	Template Instructionalization	10K
GroundingGPT's SFT (Li et al., 2024b)	Flickr30K, VCR, Activitynet Captions, Clotho...	Uni-Modal Understanding	Image,Video,Audio,Text	GPT Instructionalization	1M
AnyGPT's SFT (Du et al., 2018)	–	Uni-Modal Understanding,Cross-Modal Understanding, Uni-Modal Generation,Cross-Modal Generation	Image,Audio,Text	GPT Instructionalization, T2X generation	208K
CAT's SFT (Ye et al., 2024b)	VGGSound, AVQA, VideoInstruct100K...	Cross-Modal Understanding	Video,Audio,Text	GPT Instructionalization	100K
AVicuna's SFT (Tang et al., 2024a)	UnAV-100, VideoInstruct100K, ActivityNet Captions, DiDeMo	Uni-Modal Understanding,Cross-Modal Understanding	Video,Audio,Text	Template Instructionalization	49K
M3DBench'SFT (Li et al., 2024b)	ScanNet, ScanRefer, ShapeNet...	Uni-Modal Understanding,Cross-Modal Understanding	Image,3D,Text	Template Instructionalization,./GPT Instructionalization	320k
Uni-Moe's SFT (Li et al., 2024f)	LLaVA-Instruct-150K, LibriSpeech, VideoInstruct100K...	Uni-Modal Understanding,Cross-Modal Understanding	Image,Video,Audio,Text	Template Instructionalization, T2X generation	874K
X-VILA's SFT (Ye et al., 2024a)	WebVid, ActivityNetCaption, LLaVA-Instruct-150K...	Uni-Modal Understanding,Cross-Modal Understanding, Uni-Modal Generation,Cross-Modal Generation	Image,Video,Audio,Text	Template Instructionalization	–
EMOVA's SFT (Chen et al., 2024c)	ShareGPT-4o, MSCOCO, LLaVA-Instruct-150K...	Uni-Modal Understanding,Cross-Modal Understanding, Uni-Modal Generation,Cross-Modal Generation	Image,Audio,Text	Template Instructionalization, GPT Instructionalization, T2X generation	4.4M
VideoLLaMA2's SFT (Cheng et al., 2024b)	AVQA, AVSD, MusicCaps...	Uni-Modal Understanding,Cross-Modal Understanding	Image,Video,Text	Template Instructionalization	1.5M
PathWeave's SFT (Yu et al., 2024a)	VQAv2, MSRVTT, Cap3D...	Uni-Modal Understanding	Image,Video,Audio,3D,Depth	Template Instructionalization, T2X generation	23.2M
Spider's SFT (Lai et al., 2024)	AudioCap, CC3M, Webvid...	Cross-Modal Understanding,Cross-Modal Generation	Image,Video,Audio,Text	Template Instructionalization, GPT Generation	–
GMA's SFT (Szot et al., 2024b)	Meta-World,CALVIN, Maniskill...	Uni-Modal Understanding,Cross-Modal Understanding, Cross-Modal Generation	Image,Text,Action	Template Instructionalization	2.2M
OCTAVIUS's SFT (Chen et al., 2024g)	MSCOCO,Bamboo, ScanNet...	Uni-Modal Understanding	Image,3D,Text	Template Instructionalization, GPT Generation	–
Lyra's SFT (Zhong et al., 2024)	Mini-Gemini, Collected Youtube's Audio	Uni-Modal Understanding,Cross-Modal Understanding, Uni-Modal Generation	Image,Audio,Text	GPT Generation, T2X Generation	1.5M
video-SALMONN's SFT (Sun et al., 2024a)	LibriSpeech,AudioCaps, LLaVA-Instruct-150K...	Uni-Modal Understanding,Cross-Modal Understanding	Video,Audio,Text	Template Instructionalization, T2X Generation	–
Meerkat's SFT (Chowdhury et al., 2024)	VGG-SS,AVSIBench, AVQA,MUSIC-AVQA...	Cross-Modal Understanding	Video,Audio,Text	Template Instructionalization, GPT Generation	3M
VITA's SFT (Fu et al., 2024b)	ShareGPT4V,LLaVA-Instruct-150K, ShareGTP4o,ShareGemini...	Uni-Modal Understanding,Cross-Modal Understanding	Image,Video,Audio,Text	T2X Generation	–
Baichuan-omni's SFT (Li et al., 2024c)	vFLAN,VideoInstruct100K...	Uni-Modal Understanding,Cross-Modal Understanding	Image,Video,Audio,Text	T2X Generation	–
LongVALE-LLM's SFT (Geng et al., 2024)	LongVALE	Uni-Modal Understanding,Cross-Modal Understanding	Video,Audio,Text	GPT Generation	25.4K
UnifiedMLLM's SFT (Li et al., 2024g)	LISA,SmartEdit...	Uni-Modal Understanding,Cross-Modal Understanding, Uni-Modal Generation,Cross-Modal Generation	Image,Video,Audio,Text	Template Instructionalization, GPT Generation	100K
Dolphin's SFT (Anonymous, 2024)	AVQA,Flickr-SoundNet, VGGSound,LLP...	Uni-Modal Understanding,Cross-Modal Understanding	Video,Audio,Text	Template Instructionalization, GPT Generation	–

Table 3: The statistics for OmniMLLM’s Instruction Data, including the data sources, interaction forms, involved modalities, and construction methods.

Name	Capability Category	Modality	Specific-Task	Metrics
VQA v2 (Goyal et al., 2017)	Unimodal Understanding	Image,Text	QA	Acc
GQA (Hudson and Manning, 2019)	Unimodal Understanding	Image,Text	QA	Acc
DocVQA (Mathew et al., 2021)	Unimodal Understanding	Image,Text	QA	Acc
IconQA (Lu et al., 2021)	Unimodal Understanding	Image,Text	QA	Acc
OCR-VQA (Mishra et al., 2019)	Unimodal Understanding	Image,Text	QA	Acc
STVQA (Bilen et al., 2019)	Unimodal Understanding	Image,Text	QA	Acc
VSR (Liu et al., 2023a)	Unimodal Understanding	Image,Text	QA	Acc
Hateful Meme (Kiehl et al., 2020)	Unimodal Understanding	Image,Text	QA	AUC
OKVQA (Marino et al., 2019)	Unimodal Understanding	Image,Text	QA	Acc
VizWiz (Gurari et al., 2018)	Unimodal Understanding	Image,Text	QA	Acc
TextVQA (Singh et al., 2019)	Unimodal Understanding	Image,Text	QA	Acc
nocap (Agrawal et al., 2019)	Unimodal Understanding	Image,Text	Caption	CIDER
ScienceQA (Lu et al., 2022)	Unimodal Understanding	Image,Text	QA	Acc
MSCOCO Caption (Lin et al., 2014)	Unimodal Understanding	Image,Text	Caption	CIDER,BLEU
Flicker Caption (Plummer et al., 2015)	Unimodal Understanding	Image,Text	Caption	CIDER
Visual Dialog (Das et al., 2017)	Unimodal Understanding	Image,Text	Dialogue	MRR
RefCOCO (Yu et al., 2016)	Unimodal Understanding	Image,Text	Grounding	Acc
RefCOCO+ (Yu et al., 2016)	Unimodal Understanding	Image,Text	Grounding	Acc
RefCOCOg (Mao et al., 2016)	Unimodal Understanding	Image,Text	Grounding	Acc
A-okvqa (Schwenk et al., 2022)	Unimodal Understanding	Image,Text	QA	Acc
POPE (Li et al., 2023)	Unimodal Understanding	Image,Text	Hallucination	Acc
IIT5K (Mishra et al., 2012)	Unimodal Understanding	Image,Text	OCR	WAC(word ACC)
IC13 (Karatzas et al., 2013)	Unimodal Understanding	Image,Text	OCR	WAC(word ACC)
IC15 (Karatzas et al., 2015)	Unimodal Understanding	Image,Text	OCR	WAC(word ACC)
Total-Text (Chng and Chan, 2017)	Unimodal Understanding	Image,Text	OCR	WAC(word ACC)
CUTE80 (Rissanawan et al., 2014)	Unimodal Understanding	Image,Text	OCR	WAC(word ACC)
SVT (Wang et al., 2011)	Unimodal Understanding	Image,Text	OCR	WAC(word ACC)
SVTP (Phan et al., 2013)	Unimodal Understanding	Image,Text	OCR	WAC(word ACC)
COCO-Text (Veit et al., 2016)	Unimodal Understanding	Image,Text	OCR	WAC(word ACC)
MMB (Li et al., 2024a)	Unimodal Understanding	Image,Text	Comprehensive Benchmark	GPT ACC
MMB (Fu et al., 2023)	Unimodal Understanding	Image,Text	Comprehensive Benchmark	GPT ACC
LLaVA-Bench (Liu et al., 2023c)	Unimodal Understanding	Image,Text	Comprehensive Benchmark	GPT ACC
Mmmu (Yue et al., 2024)	Unimodal Understanding	Image,Text	Comprehensive Benchmark	GPT ACC
SEED (Se et al., 2023)	Unimodal Understanding	Image,Text	Comprehensive Benchmark	GPT ACC
MM-Vet (Yu et al., 2024d)	Unimodal Understanding	Image,Text	Comprehensive Benchmark	GPT ACC
ActivityNet-QA (Yu et al., 2019)	Unimodal Understanding	Video,Text	QA	Acc
MSRVTT-QA (Xu et al., 2016)	Unimodal Understanding	Video,Text	QA	Acc
MSVD-QA (Xu et al., 2017)	Unimodal Understanding	Video,Text	QA	Acc
HowQA (Li et al., 2020)	Unimodal Understanding	Video,Text	QA	Acc
NESTQA (Xiao et al., 2021)	Unimodal Understanding	Video,Text	QA	ACC
STAR (Wu et al., 2021)	Unimodal Understanding	Video,Text	QA	Acc
MSVD-Caption (Xu et al., 2017)	Unimodal Understanding	Video,Text	QA	CIDER
VATEX (Wang et al., 2019)	Unimodal Understanding	Video,Text	Caption	CIDER
MSRVTT-Caption (Xu et al., 2016)	Unimodal Understanding	Video,Text	Caption	CIDER,BLEU
Video-ChatGPT Benchmark (Maz et al., 2024)	Unimodal Understanding	Video,Text	Comprehensive Benchmark	GPT ACC,GPT Score
Kinetics-400	Unimodal Understanding	Video,Text	Classification	Acc
Perception test (Patrascu et al., 2023)	Unimodal Understanding	Video,Text	Comprehensive Benchmark	GPT ACC
EgoSchema (Mangalam et al., 2023)	Unimodal Understanding	Video,Text	Comprehensive Benchmark	GPT ACC
MVBench (Li et al., 2024a)	Unimodal Understanding	Video,Text	Comprehensive Benchmark	GPT ACC
VideoMME (Fu et al., 2024a)	Unimodal Understanding	Video,Text	Comprehensive Benchmark	GPT ACC
Charades-STA (Sigurdsson et al., 2016)	Unimodal Understanding	Video,Text	Grounding	IoU
AudioCaps (Kim et al., 2019)	Unimodal Understanding	Audio,Text	Caption	CIDER,SPICE,METEOR,BLEU,SPIDER
ClothoQA (Lipping et al., 2022)	Unimodal Understanding	Audio,Text	QA	Acc
VocalSound (Gong et al., 2022)	Unimodal Understanding	Audio,Text	QA	Acc
Clotho v1 (Drossos et al., 2020)	Unimodal Understanding	Audio,Text	Caption	CIDER
Clotho v2 (Drossos et al., 2020)	Unimodal Understanding	Audio,Text	Caption	CIDER
ESC50 (Piczak, 2015)	Unimodal Understanding	Audio,Text	Classification	Acc
LibriSpeech (Pangniov et al., 2015)	Unimodal Understanding	Audio,Text	ASR	WER
ASHELL-2 (Du et al., 2018)	Unimodal Understanding	Audio,Text	ASR	WER
Wenetspeech (Zhang et al., 2022a)	Unimodal Understanding	Audio,Text	ASR	WER
MusicCap (Agostinelli et al., 2023)	Unimodal Understanding	Audio,Text	Caption	CLAP Score
TUT2017 (Mesaros et al., 2016)	Unimodal Understanding	Audio,Text	Classification	Acc
EHSI (Li et al., 2024)	Unimodal Understanding	Audio,Text	QA	Acc
Cap3D Caption (Luo et al., 2023b)	Unimodal Understanding	3D,Text	Caption	CIDER
Objaverse Caption (Deitke et al., 2023)	Unimodal Understanding	3D,Text	Caption	METEOR,ROUGE,BLEU
Cap3D QA (Luo et al., 2023b)	Unimodal Understanding	3D,Text	QA	Acc
Objaverse Classification (Deitke et al., 2023)	Unimodal Understanding	3D,Text	Classification	GPT ACC
ModelNet40 (Wu et al., 2015)	Unimodal Understanding	3D,Text	Classification	Acc
ScanRefer (Chen et al., 2020a)	Unimodal Understanding	3D,Text	Grounding	mAP
N-3D (Achlioptas et al., 2020)	Unimodal Understanding	3D,Text	Caption	BLEU,CIDER,METEOR,ROUGE-L
SQA3D (Ma et al., 2023)	Unimodal Understanding	3D,Text	QA	Acc
ScanQA (Azuma et al., 2022)	Unimodal Understanding	3D,Text	QA	Acc
SUN RGB-D (Song et al., 2015)	Unimodal Understanding	Depth,Text	Classification	Acc
NYUv2 (Silberman et al., 2012)	Unimodal Understanding	Depth,Text	Classification	Acc
SUN RGB-D, generated Normal (Han et al., 2024a)	Unimodal Understanding	Normal,Text	Classification	Acc
NYUv2, generated Normal (Han et al., 2024a)	Unimodal Understanding	Normal,Text	Classification	Acc
ThermalQA (Yu et al., 2024c)	Unimodal Understanding	Thermal,Text	QA	Acc
TechQA (Yu et al., 2024c)	Unimodal Understanding	Touch,Text	QA	Acc
Ego4D (Grimman et al., 2022)	Unimodal Understanding	IMU,Text	Caption	CIDER,ROUGE
NSD (Allen et al., 2022)	Unimodal Understanding	fMRI,Depth Map,Text	Caption	CIDER,ROUGE
MSCOCO (Lin et al., 2014)	Unimodal Generation	Image,Text	TX2X Edit	FID,CLIPSIM
MSRVTT (Xu et al., 2016)	Unimodal Generation	Video,Text	TX2X Generate	CLIPSIM
AudioCaps (Kim et al., 2019)	Unimodal Generation	Audio,Text	TX2X Generate, TX2X Edit	FAD
DAVIS (Perazzi et al., 2016)	Unimodal Generation	Video,Text	TX2X Edit	CLIPSIM
UCF-101 (Soomro et al., 2012)	Unimodal Generation	Video,Text	TX2X Generate	FID,FVD,JIS,CLIPSIM
Evakraftr (Liu et al., 2024c)	Unimodal Generation	Video,Text	TX2X Generate	FVD,CLIPSIM
VCTR (Vean et al., 2017)	Unimodal Generation	Audio,Text	TX2X Generate, TX2X Edit	WER,MCD
MusicCap (Agostinelli et al., 2023)	Unimodal Generation	Audio,Text	TX2X Generate	FAD
Dreambench (Ruiz et al., 2023)	Unimodal Generation	Image,Text	TX2X Generate	CLIP-I,CLIP-T,DINO
MUSIC-AVQA (Li et al., 2022b)	Crossmodal Understanding	Video,Audio,Text	QA	Acc
AVSD (Al-Ammi et al., 2018)	Crossmodal Understanding	Video,Audio,Text	Dialogue	CIDER,BLEU
RACE-Audio (Li et al., 2024f)	Crossmodal Understanding	Image,Audio,Text	Comprehensive Benchmark	Acc
VALOR Caption (Chen et al., 2023c)	Crossmodal Understanding	Video,Audio,Text	Caption	CIDER,BLEU
MMBench-Audio (Li et al., 2024f)	Crossmodal Understanding	Image,Audio,Text	Comprehensive Benchmark	Acc
AVQA (Li et al., 2022a)	Crossmodal Understanding	Video,Audio,Text	QA	Acc
MCUB (Chen et al., 2024a)	Crossmodal Understanding	Image,Video,Audio,3D,Text	Comprehensive Benchmark	Acc
DisCn (Panagopoulou et al., 2024)	Crossmodal Understanding	Image,Video,Audio,3D	Comprehensive Benchmark	Acc
OmniRn (Chen et al., 2024d)	Crossmodal Understanding	Image,Video,Audio,Text	Comprehensive Benchmark	Acc
Curv (Leng et al., 2024)	Crossmodal Understanding	Image,Video,Audio,Text	Hallucination	Acc
ISQA (Sun et al., 2023a)	Crossmodal Understanding	Image,Audio,Text	QA	Acc
VGSound (Chen et al., 2020b)	Crossmodal Understanding	Video,Audio,Text	QA	Acc
VATEX (Wang et al., 2019)	Crossmodal Understanding	Video,Audio,Text	Caption	CIDER
UnAV-100 (Gong et al., 2023)	Crossmodal Understanding	Video,Audio,Text	Ground	IoU
LLP (Tian et al., 2020)	Crossmodal Understanding	Video,Audio,Text	Ground	IoU
PresentationQA (Sun et al., 2024a)	Crossmodal Understanding	Video,Audio,Text	QA	ACC
LogVALE Caption	Crossmodal Understanding	Video,Audio,Text	QA	CIDER
TVL Benchmark (Fu et al., 2024c)	Cross-modal Understanding	Touch,Image,Text	Caption	ACC
AVEB (Sun et al., 2023a)	Crossmodal Understanding,Unimodal Understanding	Image,Video,Audio,Text	Comprehensive Benchmark	ACC,METEOR,SPIDER,WER
XioX Benchmark (Ye et al., 2024a)	Crossmodal Understanding,Crossmodal Generation	Image,Video,Audio,Text	Comprehensive Benchmark	X-to-X Alignment Score

Table 4: An overview of benchmarks and tasks of Omni-MLLMs, including the abilities being evaluated, the involved modalities, specific tasks, and evaluation metrics.

Model	LLM	Uni-Modal Understanding								Uni-Modal Generation			Cross-Modal Understanding			
		MSVD-QA	MSRVTT-QA	VQA ^{v2 test}	Flicker	MMB ^{cm}	AudioCaps ^{cm test}	ClothoAQA	Objaverse	COCO ^{gen}	AudioCaps ^{gen}	MSRVTT ^{gen}	VGGSS	AVSD	MUSIC-AVQA	AVQA
Omni-MLLMs																
eP-ALM	OPT-2.7B	38.4	38.51	54.47	—	—	61.86	—	—	—	—	—	—	—	—	—
ChatBridge 13B	Vicuna-13B	45.3	—	—	82.5	—	—	—	—	—	—	—	—	—	43	—
PandaGPT	Vicuna-13B	46.7	23.7	—	—	—	—	—	—	—	—	—	32.7	26.1	33.7	79.8
Video-LLaMA	Vicuna-7B	51.6	29.6	—	—	—	—	—	—	—	—	—	40.8	36.7	36.6	81
Macaw-LLM	LLaMA-7B	42.1	25.5	—	—	3.84	33.3	—	—	—	—	—	36.1	34.3	31.8	78.7
ImageBind-LLM	LLaMA-7B	—	—	—	23.49	—	—	10.3	31	—	—	—	—	—	39.72	54.26
NEXT-GPT	Vicuna-7B	64.5	58.4	66.7	84.5	58	81.3	—	—	10.07	8.67	31.97	—	—	—	—
AnyMAL 13B	LLaMA2-13B	—	—	59.6	—	—	—	—	—	—	—	—	—	—	—	—
AnyMAL 70B	LLaMA2-70B	—	—	64.2	95.9	—	77.8	—	—	—	—	—	—	—	—	—
X-InstructBLIP 7B	Vicuna-7B	51.7	41.3	30.61	82.1	8.96	67.9	15.4	50	—	—	—	—	—	28.1	—
X-InstructBLIP 13B	Vicuna-13B	49.2	—	—	74.7	—	53.7	21.7	—	—	—	—	20.3	52.1	44.5	44.23
OneLLM 7B	LLaMA2-7B	56.5	—	71.6	78.6	60	—	57.9	44.5	—	—	—	—	—	47.6	—
AV-LLM	Vicuna-7B	67.3	53.7	—	—	—	35.5	—	—	—	—	—	47.6	52.6	45.2	—
UIO-2xcl 6.8B	—	52.2	41.5	79.4	—	71.5	48.9	—	—	13.39	5.89	—	—	—	—	—
ModaVerse	Vicuna-7b	—	56.5	—	—	—	79.2	—	—	11.24	8.22	30.14	—	—	—	—
CREMA 7B	Mistral-7B	—	—	—	—	—	—	—	—	—	—	—	—	—	52.6	—
GroundingGPT	Vicuna-7B	67.8	51.6	78.7	—	63.8	—	—	—	—	—	—	—	—	—	—
NaiveMC	Vicuna-7B	—	—	—	—	—	—	—	55	—	—	—	—	—	53.63	80.7
DAMC	Vicuna-7B	—	—	—	—	—	—	60.5	—	—	—	—	—	—	57.32	81.31
AnyGPT	LLaMA2-7B	—	—	—	—	—	—	—	—	—	—	—	—	—	—	—
CAT	LLaMA2-7B	—	62.7	—	—	—	—	—	—	—	—	—	—	48.6	92	—
AVicuna	Vicuna-7B	70.2	59.7	—	—	—	—	—	—	—	—	—	—	53.1	49.6	—
Uni-MoE	LLaMA-7B	55.6	—	66.2	—	69.82	—	32.6	—	—	—	—	—	—	—	—
X-VILA 7B	Vicuna-7B	—	—	72.9	—	—	—	—	—	—	—	—	—	—	—	—
VideoLLaMA2-7B	Mistral-7B	71.7	—	—	—	—	—	—	—	—	—	—	71.4	57.2	80.9	—
Meerkat	Llama-2-7B-Chat	—	—	—	—	—	—	—	—	—	—	—	—	—	—	87.14
InterOmni	InterLM-2-Chat-7B	—	—	—	—	81.7	—	—	—	—	—	—	—	—	—	—
UnifiedMLLM	Vicuna-7B	—	—	—	—	—	—	—	—	—	—	—	—	—	—	—
VITA	Mixtral-8x7B	—	—	—	—	71.8	—	—	—	—	—	—	—	—	—	—
EMOVA	LLaMA-3.1-8B	—	—	—	—	82.8	—	—	—	—	—	—	—	—	—	—
BaiChuan-omni-7B	—	72.2	—	—	—	76.2	—	—	—	—	—	—	—	—	—	—
OMCAT	Vicuna-7B	—	—	—	—	—	—	—	—	—	—	—	—	49.4	73.8	90.2
PathWeave-7B	Vicuna-7B	47.8	37.4	—	—	64	33.5	—	—	—	—	—	—	—	—	—
Spider	Llama-2-7B	—	—	—	—	81.7	—	—	—	11.23	8.18	30.97	—	—	—	—
Specife-MLLMs																
VILA-7B	LLaMA-2-7B	—	—	79.9	74.7	68.9	—	—	—	—	—	—	—	—	—	—
VideoChat2	Vicuna-7B	70	54.1	—	—	—	—	—	—	—	—	—	—	—	—	—
Qwen-Audio	Qwen-7B	—	—	—	—	—	—	57.9	—	—	—	—	—	—	—	—
PointLLM	Vicuna-7B	—	—	—	—	—	—	—	47.5	—	—	—	—	—	—	—
Emu-13B	—	—	—	52	—	—	—	—	—	11.66	—	—	—	—	—	—
Video-LaVIT	Llama2-7B	73.2	—	80.3	—	67.3	—	—	—	—	—	30.12	—	—	—	—

Table 5: **The performance of Omni-MLLMs on different benchmarks.** The selected uni-modal understanding benchmarks include Video-Text2Text (Xu et al., 2017, 2016), Image-Text2Text (Goyal et al., 2017; Plummer et al., 2015; Liu et al., 2024d), Audio-Text2Text (Kim et al., 2019; Lipping et al., 2022), and 3D-Text2Text (Deitke et al., 2023). The chosen uni-modal generation benchmarks include Text2Image (Lin et al., 2014), Text2Video (Xu et al., 2016), and Text2Audio (Kim et al., 2019). The selected cross-modal understanding benchmarks are Image-Audio-Text2Text (Chen et al., 2020b; AlAmri et al., 2018) and Video-Audio-Text2Text (Li et al., 2022b,a).