

AV-ODYSSEY BENCH: FROM FUNDAMENTAL AUDIO PERCEPTION TO AUDIO-VISUAL UNDERSTANDING

Anonymous authors

Paper under double-blind review

ABSTRACT

Recent multimodal large language models (MLLMs), such as GPT-4o, Gemini 1.5/2.5 Pro, and Reka Core, have advanced audio-visual reasoning capabilities, achieving strong performance in tasks like cross-modal understanding and generation. However, our **DeafTest** uncovers unanticipated failures: most of the state-of-the-art MLLMs struggle with very simple audio tasks, such as *distinguishing louder sounds* or *sound counting*. This raises a fundamental question—does a deficiency in low-level audio perception constrain higher-level audio-visual reasoning? To address this, we introduce **AV-Odyssey Bench**—a comprehensive benchmark of 4,555 meticulously designed problems that integrate text, audio, and visual modalities. Each task requires models to unify cross-modal reasoning, leveraging synchronized audio-visual cues to infer solutions. By structuring questions as multiple-choice, we ensure objective, reproducible evaluations without reliance on subjective human or LLM-based judgments. Through comprehensive benchmarking of closed-source and open-source models, we showcase: (i) current MLLMs lack robust audio-visual integration ability and (ii) performance on DeafTest (Pearson’s $r = 0.945$) strongly correlates with AV-Odyssey accuracy. These findings not only challenge prevailing assumptions about the “multimodal proficiency” of leading models, but also highlight the importance of fundamental audio perception as a bottleneck for audio-visual reasoning. We believe that our results provide concrete guidance for future research in future dataset design, alignment strategies, and architectures toward truly integrated audio-visual understanding.

1 INTRODUCTION

Multimodal reasoning has advanced significantly through two key stages: vision-language models (VLMs) and their evolution into audio-visual extensions. Early VLMs, such as GPT-4V(ision) (OpenAI, 2023), pioneered visual perception capabilities, enabling tasks like object counting (Xu et al., 2023), numerical reasoning on tabular data (Yang et al., 2023), and geometric problem-solving (Zhang et al., 2025). Building on this foundation, modern Multimodal Large Language Models (MLLMs)¹ integrate *audio* modalities, exemplified by GPT-4o (Hurst et al., 2024), Gemini 1.5 (Team et al., 2024a), and Gemini 2.5 Pro (Google, 2025). These models push the boundaries of multimodal reasoning, achieving strong performance in automatic speech recognition (ASR) (Hurst et al., 2024), cross-modal translation (Team et al., 2024a), and audio-visual captioning (Han et al., 2024; Zhan et al., 2024).

Benchmarking is a critical component of the community, as it helps specify the development direction. Prior work mainly focus on visual problem-solving, such as general comprehension (Li et al., 2024b; Liu et al., 2023c; Fu et al., 2023) and mathematical reasoning (Chen et al., 2021; Cao & Xiao, 2022; Chen et al., 2022; Zhang et al., 2025; Lu et al., 2023). On the other hand, audio-visual benchmarks such as AVQA (Yang et al., 2022), OmniBench (Li et al., 2025), and MusicAVQA (Li et al., 2022) focus on testing MLLMs with audio-visual tasks that require the simultaneous processing and integration of visual and auditory information.

In this paper, we directly test MLLMs’ ability to *see* instead of reasoning on simple low-level audio task. This is inspired by the BlindTest (Rahmanzadehgervi et al., 2024), which reveals that powerful

¹In this work, MLLMs specifically refer to audio-vision LLMs, distinct from vision-only VLMs.

vision language models are still struggling with very simple vision tasks that are easy for humans. Concretely, we propose a DeafTest benchmark (Fig. 1), including four extremely simple audio tasks inspired from Schwabach test (Huizing, 1975). We test a set of powerful MLLMs like Gemini 1.5 (Team et al., 2024a), Gemini 2.5 Pro (Google, 2025), Reka (Team et al., 2024b), and GPT-4o (Hurst et al., 2024) on these easy task that only involve basic sound properties (*e.g.*, loudness, pitch, duration), as shown in Table 1. Our key findings are:

- Despite their ability to recognize complex speech content, MLLMs do not perform as well as expected on sound counting tasks. The best-performing model, Gemini 1.5 Pro, achieves only 81%, while humans can easily score 100%. The sounds in these tasks are monotonous and are clearly separated by silent intervals within the audio clip.
- Most MLLMs appear to be insensitive to sound loudness or sound pitch, except for Gemini 2.5 Pro and Qwen-2.5-Omni-7B. Models are required to distinguish the louder sound or the higher pitch from two given sounds. Several models perform under 60%, while a random guess baseline is 50%.
- The duration comparison task presents models with two sounds and asks them to determine which has the longer duration. Model performance varies significantly, where Gemini 2.5 Pro achieves a high score of 99.0%, models such as Reka Core and Reka Edge perform poorly, with scores of 40.0% and 44.0%, respectively.

These findings are noteworthy: despite the strong performance of proprietary models on complex tasks such as automatic speech recognition, their performance on basic audio perception tasks reveals a noticeable gap compared to human capabilities. Notably, Gemini 2.5 Pro (Google, 2025) generally outperforms other models. This raises a critical question: Does a deficiency in low-level audio perception limit a model’s ability to perform higher-level audio-visual reasoning? Furthermore, can the performance on DeafTest serve as a key indicator of a model’s overall audio-visual reasoning ability?

To address this, we introduce a novel and comprehensive benchmark **AV-Odyssey** to evaluate the audio-visual reasoning performance of MLLMs. This benchmark comprises 4,555 questions across 26 tasks, spanning a wide range of audio attributes, each strategically engineered to require cross-modal synergy (Fig. 3). Key design principles include: 1) multimodal necessity: questions are filtered using VLMs and audio models to exclude tasks solvable by single-modal reasoning; 2) diverse scopes: coverage of sound attributes (timbre, spatial dynamics), application domains (music, transportation), and temporal reasoning; 3) objective evaluation: multiple-choice format eliminates reliance on subjective human or LLM-based assessments.

We benchmark closed-source (GPT-4o (Hurst et al., 2024), Gemini 1.5 (Team et al., 2024a), Gemini 2.5 Pro (Google, 2025)) and open-source models (Han et al., 2024; Lu et al., 2022; Zhan et al., 2024; Wu et al., 2023; Su et al., 2023; Cheng et al., 2024; Fu et al., 2024) on our proposed AV-Odyssey. Combined with DeafTest results, the key findings are:

- Overall, current MLLMs still fall short in processing complex audio-visual information integration tasks.
- Performance on DeafTest strongly correlates with AV-Odyssey accuracy with a Pearson’s $r = 0.945$ (Fig. 2). In other words, models that perform poorly on simple audio tasks (DeafTest) also consistently underperform on AV-Odyssey.
- Error analysis (see Sec. 4.3) shows that most audio-visual inference errors stem from mis-perceiving audio inputs. This finding aligns with findings on DeafTest, where models with weak performance on simple auditory tasks also underperform on AV-Odyssey.

To conclude, this work systematically investigates the audio-visual comprehension capabilities of current MLLMs through two complementary perspectives: DeafTest and AV-Odyssey Bench. **Our main contributions are as follows:** (1) DeafTest provides the **first** dedicated evaluation of MLLMs’ basic listening abilities (*e.g.*, pitch and loudness discrimination), uncovering critical weaknesses in fundamental audio perception. (2) AV-Odyssey Bench introduces a large-scale and comprehensive benchmark for advanced cross-modal reasoning, spanning diverse audio attributes, visual contexts, and real-world scenarios. (3) We empirically demonstrate a strong correlation between fundamental audio perception and overall audio-visual reasoning performance, highlighting the dependency

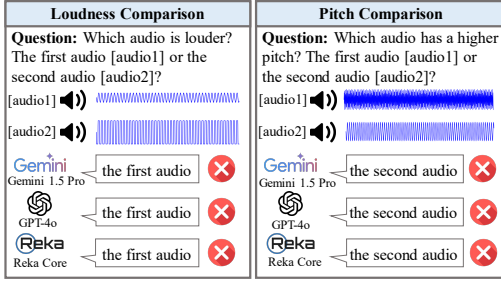


Figure 1: Illustration of two out of four Deaf-Test tasks: loudness and pitch comparison.

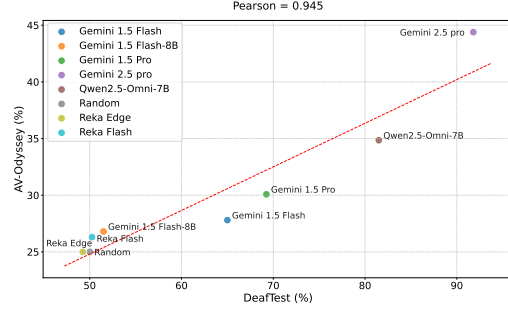


Figure 2: Performance on DeafTest and AV-Odyssey, showing a strong positive correlation.



Figure 3: Overview of AV-Odyssey Benchmark. AV-Odyssey Bench demonstrates three major features: 1. Comprehensive Audio Attributes; 2. Extensive Domains; 3. Interleaved Text, Audio, and Images.

of higher-level multimodal reasoning on low-level perceptual skills. By explicitly linking low-level perception to high-level reasoning, our benchmarks offer the first comprehensive diagnosis of MLLMs' multimodal capabilities. We reveal that when pursuing strong audio-visual comprehension ability, it's noteworthy to monitor the model performance on fundamental audio perception performance. These resources provide a foundation for targeted advances in dataset construction, alignment strategies, and model architectures, ultimately paving the way toward truly integrated audio-visual intelligence.

2 RELATED WORK

Multimodal Large Language Models. Large language models (LLMs) have demonstrated remarkable performance across diverse textual domains (OpenAI, 2023; Radford, 2018; Brown, 2020; Touvron et al., 2023; Bi et al., 2024). The success of these models has catalyzed significant advancements in vision language models and multimodal large language models. Inspired by the textual prowess of LLMs, vision language models have emerged to extend computational capabilities into visual comprehension. These models enable LLMs to perform sophisticated visual tasks, including visual question answering (Liu et al., 2023a; Li et al., 2023a; Zhu et al., 2023; Liu et al., 2023b; Ye et al., 2023; Dai et al., 2023; Bai et al., 2023; Zhang et al., 2023b), visual grounding (Peng et al., 2023; Wang et al., 2023; Chen et al., 2023a;b), document understanding (Ye et al., 2023; Hu et al., 2024; Zhang et al., 2023c; Lv et al., 2023), long video understanding (Liu et al., 2024; Shen et al., 2024; Zhang et al., 2024a; Li et al., 2024c; Ren et al., 2024). Building upon vision-language achievements, researchers have further expanded multimodal horizons by integrating the audio modality (Han et al., 2024; Lu et al., 2022; Zhan et al., 2024; Wu et al., 2023; Su et al., 2023; Fu et al., 2024; Cheng et al., 2024; Chowdhury et al., 2024; Xu et al., 2025; Lu et al., 2025). These advanced models now accommodate audio inputs, further expanding the landscape of multimodal artificial intelligence.

Benchmarking Multimodal Large Language Models. The rapid development of vision language models has been accompanied by the emergence of specialized benchmarks to assess their perfor-

Table 1: Results on four basic auditory tasks (DeafTest). The questions are designed as two-choice questions. The random baseline performance is 50%. The final column shows the average performance across the four tasks.

Method	Sound Counting	Loudness Comparison	Pitch Comparison	Duration Comparison	Average
Random	50.0	50.0	50.0	50.0	50.0
Gemini 1.5 Flash (Team et al., 2024a)	55.0	62.0	54.0	89.0	65.0
Gemini 1.5 Flash-8B (Team et al., 2024a)	49.0	55.0	51.0	51.0	51.5
Gemini 1.5 Pro (Team et al., 2024a)	81.0	60.0	52.0	84.0	69.3
Gemini 2.5 Pro (Google, 2025)	69.0	100.0	98.0	99.0	91.5
Reka Core (Team et al., 2024b)	54.0	43.0	42.0	40.0	44.8
Reka Flash (Team et al., 2024b)	48.0	58.0	51.0	44.0	50.3
Reka Edge (Team et al., 2024b)	47.0	56.0	50.0	44.0	49.3
GPT-4o audio-preview (Hurst et al., 2024)	50.0	58.0	58.0	57.0	55.8
Qwen-2.5-Omni-7B (Xu et al., 2025)	60.0	100.0	86.0	80.0	81.5

Table 2: Comparisons between MLLM benchmarks / datasets.

Benchmark / Dataset	Modality	Questions	Answer Type	Customized Question	Audio Attributes						Multiple Domains		Interleaved
					Timbre	Tone	Melody	Space	Time	Hallucination	Intricacy	✓	
MME Bench (Fu et al., 2023)	Image	2194	Y/N	✓	-	-	-	-	-	-	-	✓	✗
MMBench (Liu et al., 2023c)	Image(s)	2974	A/B/C/D	✓	-	-	-	-	-	-	-	✓	✗
SEED-Bench-2 (Li et al., 2024b)	Image(s) & Video	24371	A/B/C/D	✓	-	-	-	-	-	-	-	✓	✓
AVQA Dataset (Yang et al., 2022)	Video & Audio	57335	A/B/C/D	✓	✓	✓	✗	✗	✗	✗	✓	✓	✗
Pano-AVQA Dataset (Yun et al., 2021)	Video & Audio	51700	defined words & bbox	✓	✓	✓	✗	✓	✓	✗	✓	✓	✗
Music-AVQA Dataset (Li et al., 2022)	Video & Audio	45867	defined words	✓	✓	✓	✗	✓	✓	✓	✓	✓	✗
SAVE Bench (Sun et al., 2024)	Image & Video & Audio	4350	free-form	✗	✓	✓	✗	✗	✗	✗	✓	✓	✗
OmniBench (Li et al., 2025)	Image & Audio	1142	A/B/C/D	✓	✓	✓	✗	✗	✗	✗	✓	✓	✗
AV-Odyssey Bench (ours)	Image(s) & Video & Audio(s)	4555	A/B/C/D	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓

mance across various domains (Yue et al., 2024; Fu et al., 2023; Li et al., 2024b; Chen et al., 2021; Lu et al., 2023). A significant subset of these benchmarks focuses on vision comprehension (Yue et al., 2024; Fu et al., 2023; Li et al., 2024b;a) and mathematical reasoning capabilities (Chen et al., 2021; Cao & Xiao, 2022; Lu et al., 2023; Zhang et al., 2025; Yue et al., 2024; Seo et al., 2015). However, current audio-visual benchmarks (Yang et al., 2022; Li et al., 2022; Yun et al., 2021; Sun et al., 2024; Li et al., 2025; Leng et al., 2024; Sung-Bin et al., 2024; Tang et al., 2024) face significant limitations in comprehensively assessing multimodal large language models (MLLMs). First, they predominantly focus on high-level visual tasks and neglect to explore the basic auditory perception limitations. Secondly, they do not comprehensively evaluate all attributes of the audio, comparison is detailed in Table 2. This paper begins by evaluating basic audio tasks to highlight shortcomings in auditory perception and introduces the AV-Odyssey benchmark, covering diverse audio attributes and domains. By leveraging both evaluations, we reveal key limitations of existing models and demonstrate a strong correlation between low-level audio perception and higher-level cross-modal reasoning. We emphasize that effective audio-visual comprehension requires attention to fundamental audio perception, a factor not emphasized by previous work.

3 METHOD

3.1 DEAFTEST TASKS

Drawing inspiration from the Schwabach test (Huizing, 1975), we introduce DeafTest, a suite of four extremely simple auditory tasks that critically examine the fundamental audio perception capabilities of Multi-modal Large Language Models (MLLMs). DeafTest includes the determination of the number of sounds, identification of the louder sound, recognition of the sound with a higher pitch, and detection of the sound with a longer duration. We hypothesize that MLLMs may not perform as well as expected on these basic tasks. This potential shortcoming arises from the training objectives of these models, which primarily focus on achieving high-level semantic alignment between different modalities. Consequently, this approach tends to overlook the effective utilization of low-level auditory information, which is crucial for accurately processing and understanding basic sound characteristics.

1. Count the Number of Sounds. This task evaluates MLLMs’ ability to count distinct sounds in audio clips containing 3 to 8 monotonous, clearly separated sounds. Despite strong ASR performance (e.g., GPT-4o’s 3% word error rate (OpenAI)), this tests basic auditory segmentation. Each trial presents a two-choice question. We curate 100 questions in total.

Table 3: Detailed statistics of the AV-Odyssey Benchmark.

Statistics	Number
Total Questions	4555
Total Tasks	26
Domains	10
Multi-Image, Single-Audio	2610
Single-Image, Multi-Audio	891
Single-Image, Single-Audio	434
Single-Video, Single-Audio	220
Single-Video, Multi-Audio	400
Correct Option Dist. (A:B:C:D)	1167:1153:1119:1116
Avg. Audio Time (s)	16.32
Avg. Image Res. (px)	1267×891
Avg. Video Res. (px)	1678×948
Avg. Video Time (s)	15.58

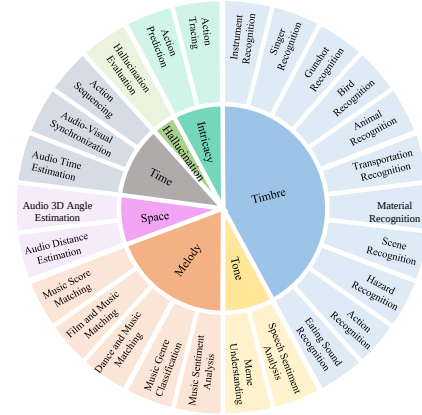


Figure 4: Overview of the 26 evaluation tasks, categorized into 7 classes based on sound attributes.

2. Discriminate the Louder Sound. In this task, we test the basic ability of MLLMs to distinguish between the loudness of sounds. The goal of MLLMs is to discriminate which sound is louder out of two given audio clips. Specifically, the decibel for quieter audio ranges from 30 dB to 60 dB, while the decibel for louder audio ranges from 70 dB to 100 dB. We randomly sample decibels from these two ranges to create two audio clips. In addition, we randomly switch the input order of the two audio clips; that is, for some questions, the quiet audio comes first, and for the rest, the loud audio comes first. Similarly, the question format is also a two-choice question.

3. Discriminate the Higher Pitch. This task focuses on pitch differentiation. Clips feature lower-pitched sounds (100–500 Hz) vs. higher-pitched ones (1000–2000 Hz), validated for human discernibility. Input order is randomized, with two-choice questions (100 total).

4. Recognize the Duration of Sound. We also test MLLMs with the duration of sound. In this task, we simplify the question by giving two audio clips of different durations. We sample the duration from 1s to 3s for the short audio, while we sample from 4s to 6s for the long audio. Similar to task 2, we provide the MLLMs with two audio clips, asking them to identify the longer one.

The DeafTest results in Table 1 reveal significant variation in model performance, with Gemini 2.5 Pro and Qwen-2.5-Omni-7B outperforming others. However, for all other models, performance falls far below expectations, with none exceeding 62%, particularly on tasks such as loudness and pitch comparison. These findings highlight substantial limitations in basic auditory perception among current MLLMs. This raises the question: could deficiencies in basic audio perception impair performance on higher-level audio-visual reasoning tasks? To investigate, we introduce the AV-Odyssey benchmark, which reveals a strong correlation between fundamental audio perception and overall audio-visual reasoning, as shown in Fig. 2.

3.2 OVERVIEW OF AV-ODYSSEY BENCH

Our AV-Odyssey Bench is a meticulously curated benchmark designed to comprehensively assess the audio-visual capabilities of MLLMs. To ensure a robust and unbiased assessment, all questions in AV-Odyssey are structured as multiple-choice, with four options per question, and options can be presented in various formats, including text, images, or audio clips. To mitigate format-specific biases, we have curated five distinct multi-choice question types. Additionally, all inputs, including text, image/video, and audio clips, are fed into MLLMs in an interleaved manner.

We compare our AV-Odyssey benchmark with previous MLLM benchmarks and datasets in Table 2. It can be found that previous works suffer from certain limitations, such as restricted audio attributes, which fail to capture the full spectrum of auditory complexity and the absence of interleaved settings, crucial for assessing real-world multimodal integration capabilities. For instance, OmniBench (Li et al., 2025) lacks multiple audio attributes, making it difficult to comprehensively assess the capabilities of MLLMs in audio-visual tasks. In contrast, our AV-Odyssey encompasses 26 tasks across 10 diverse domains and spans 7 audio attributes, with interleaved and customized

Animal Recognition Question: Out of the four audio pieces [audio1] [audio2] [audio3] [audio4], which one is the animal audio that corresponds with the image [img1]? A: the third audio B: the second audio C: the fourth audio D: the first audio [img1]  [audio1]  audio content: cow sound [audio2]  audio content: lion sound [audio3]  audio content: frog sound [audio4]  audio content: donkey sound category: timbre; domain: animals	Hazard Recognition Question: Which potentially hazardous scene shown in [img1], [img2], [img3], or [img4] corresponds best to the sound in [audio1]? A: the fourth image B: the second image C: the third image D: the first image [audio1]  audio content: burning sound [img1]  [img2]  [img3]  [img4]  category: timbre; domain: scenes	Meme Understanding Question: Based on [audio1] and [video1], what makes this meme funny? A: The meme is funny ... 'Oh My God' is delivered in a calm, monotone voice, ... facial expression remains neutral ... B: ... the exaggerated tone and the person's wide-eyed ... and the way the phrase is said adds a musical quality that ... C: The humor arises ... is yelled repeatedly in an intense voice, but the facial expression stays completely blank, while ... D: The meme is funny because ... in a cheerful, upbeat tone, ... smiling face contrasts humorously with the intense... [audio1]  audio content: "oh my god" [video1]  category: tone; domain: memes
Music Sentiment Analysis Question: Which image [img1] [img2] [img3] [img4] do you think best matches the emotion in the music [audio1]? A: the third image B: the second image C: the first image D: the fourth image [audio1]  audio content: happy music [img1]  [img2]  [img3]  [img4]  category: melody; domain: music	Dance and Music Matching Question: Please select the audio [audio1], [audio2], [audio3], or [audio4] that you think best corresponds to the dance in [video1]. A: the second audio B: the first audio C: the third audio D: the fourth audio [video1]  [audio1]  krump dance music [audio2]  breaking dance music [audio3]  waacking dance music [audio4]  LA-style hiphop dance music category: melody; domain: music	Audio Distance Estimation Question: Based on [audio1], could you provide the distance of the sound next to the woman with white clothing in [img1]? The distance should be measured egocentrically in centimeters. A: 190 B: 153 C: 182 D: 159 [audio1]  spatial audio with 4 channels [img1]  category: space; domain: daily life
Audio Time Estimation Question: Based on [audio1], what are the times for the action outlined in [video1]? A: start time: 0.00 s, end time: 3.58 s B: start time: 3.45 s, end time: 4.52 s C: start time: 4.42 s, end time: 16.79 s D: start time: 17.21 s, end time: 18.55 s [audio1]  audio content: a sequence of cleaning sounds for a cup [video1]  category: time; domain: daily life	Hallucination Evaluation Question: Which of the instruments depicted in [img1], [img2], [img3], or [img4] is not part of the sound of [audio1]? A: the third image B: the first image C: the fourth image D: the second image [audio1]  audio content: music clip with guitar, piano and drums [img1]  [img2]  [img3]  [img4]  category: hallucination; domain: music	Action Prediction Question: After [img1], what action is this person taking as heard in [audio1]? A: fill plate B: get meat mix from pan C: put meat mix on dough D: move plate [audio1]  audio content: the sound of objects scraping in the kitchen [img1]  category: intricacy; domain: daily life

Figure 5: Sampled examples from our AV-Odyssey Benchmark.

questions. The detailed statistics are shown in Table 3. This design enables an exhaustive evaluation of MLLMs, providing a nuanced and thorough assessment of their performance in complex, real-world audio-visual scenarios.

Here, we will briefly introduce the task categories that span a broad spectrum of audio attributes, including Timbre, Tone, Melody, Spatial characteristics, Temporal dynamics, and Hallucination detection. The detailed task distribution and task examples are shown in Fig. 4 and Fig. 5, respectively.

Timbre Tasks. To test the concept of matching across vision and audio modalities, MLLMs are required to match audio-visual pairs (e.g., lion’s roar sound with lion images) in timbre tasks. In addition, we have designed advanced tasks that demand internal expert-level knowledge learned from the large-scale pretraining data to solve, such as singer recognition and bird species identification.

Tone Tasks. These tasks target evaluating MLLMs with speech sentiment analysis and meme understanding. For example, meme understanding requires MLLMs to infer humorous reasons simultaneously from the voice tone and visual context.

Melody Tasks. For evaluating melody understanding abilities, we propose melody tasks. For example, the dance and music matching task requires the MLLM to understand the melody of the music and identify the one that aligns with the dance in a video.

Space Tasks. To test the spatial inference with audio and visual information, space tasks require MLLMs to infer the distance of a certain object producing a sound or to determine the 3D angle.

Time Tasks. These tasks test the cross-modal matching and temporal correlation abilities at the same time. For example, audio time estimation requires MLLMs to determine the start and end time of an action.

Hallucination Tasks. Inspired by POPE (Li et al., 2023b) that indicates severe object hallucination existing in vision language models, we designed this task to assess the hallucination issue in audio-visual reasoning.

Intricacy Tasks. These tasks challenge MLLMs to perform integrated analysis or reasoning through both visual and audio inputs, leveraging multiple attributes. For example, action prediction requires models to infer actions based on visual elements alongside various audio attributes, such as timbre and timing.

These diverse tasks provide a rigorous and multifaceted assessment of MLLMs’ audio-visual information integration capabilities, systematically probing the depth, nuance, and complexity of cross-modal perception and reasoning.

3.3 DATA CURATION PROCESS

Data Collection. AV-Odyssey Bench is an audio-visual benchmark to evaluate whether MLLMs truly have audio-visual reasoning capability. Since the audio is the newly added modality by these omni-modal models, and there is already an array of visual benchmarks, we put our attention on the attributes of sound in the benchmark construction. We first design 26 tasks that cover a wide range of audio attributes and application domains, then we manually curate each question.

Data source. Our questions involve audio, images, and videos. We use public datasets as an important source of fetching audios (MISC; AYDEMİR; Microsoft). For example, we use the instrument audios from MISC and bird sound audios from Microsoft. We mainly use images from the internet and manually filter out low-quality images. The sources of videos are either from existing datasets (Shimada et al., 2024; Damen et al., 2018) or from the internet. For example, we download meme videos from video websites to create a meme understanding question. We will detail the data source of each task in the appendix.

Quality Control. Each question-answer pair is verified by two different experts. Duplicated text information that describes visual inputs will induce MLLMs to bypass the visual input to directly derive the answer by memorizing the answer from the internet-scale training dataset (Chen et al., 2024a). Inspired by this, we first ensure that our text questions’ context is as simple as possible. Then we filter out those questions that have redundant images or audio clips by leveraging VLMs and audio LLMs. Specifically, we test all the curated questions with VLM: InternVL2 (Chen et al., 2024b), Qwen2-VL (Wang et al., 2024), MiniCPM-V 2.5 (Yao et al., 2024), BLIP3 (Xue et al., 2024), and VILA1.5 (Lin et al., 2024) and audio LLM Qwen-Audio (Chu et al., 2023), Qwen2-Audio (Chu et al., 2024), SALMONN (Tang et al., 2023), and Typhoon-Audio (Manakul et al., 2024), and filter out those questions that either of these models can solve. In experiment, 2.54% questions are filtered out because they are solved by all audio LLMs or VLMs

4 EXPERIMENT

Model Testing. We evaluate closed-source and open-source MLLMs supporting text, image/video, and audio inputs. Experiments are conducted in a **zero-shot setting** (no finetuning or few-shot prompting) to assess inherent multimodal reasoning capabilities. **Prompt Design.** Text prompts are minimized to exclude redundant information. To ensure robustness: 1. Predefine multiple question templates, randomly selecting one per evaluation. 2. Conduct three trials per question to mitigate stochasticity in model outputs.

4.1 MODEL EVALUATION AND METHODOLOGY

Evaluated Models. We assess 21 models: 9 closed-source (Gemini 2.5 Pro (Google, 2025), Gemini 1.5 Flash/Pro (Team et al., 2024a), Reka Core/Flash/Edge (Team et al., 2024b), GPT-4o (Hurst et al., 2024)) and 12 open-source (Qwen-2.5-Omni-7B (Xu et al., 2025), AV-Reasoner (Lu et al., 2025), Unified-IO-2 L/XL/XXL (Lu et al., 2022), PandaGPT (Su et al., 2023), VideoLLaMA (Zhang et al.,

Table 4: Evaluation results of various MLLMs in different parts of AV-Odyssey Bench. The highest performance is highlighted in bold. $\bar{T}$ is the averaged accuracy across corresponding dimensions, and $R_{\bar{T}}$ is the rank based on the averaged accuracy. “All Avg.” represents the averaged accuracy over all questions in our AV-Odyssey Bench.

	Model	LLM Size	Timbre		Tone		Melody		Space		Time		Hallucination		Intricacy		All Avg.	
			$\bar{T}$	$R_{\bar{T}}$	$\bar{T}$	$R_{\bar{T}}$	$\bar{T}$	$R_{\bar{T}}$	$\bar{T}$	$R_{\bar{T}}$	$\bar{T}$	$R_{\bar{T}}$	$\bar{T}$	$R_{\bar{T}}$	$\bar{T}$	$R_{\bar{T}}$	$\bar{T}$	$R_{\bar{T}}$
	Random	-	25.0	14	25.0	11	25.0	20	25.0	16	25.0	19	25.0	14	25.0	18	25.0	21
Open Source	Unified-IO-2 L (Lu et al., 2022)	1B	23.8	20	24.1	14	28.8	8	15.0	22	26.8	10	30.0	7	30.4	12	26.0	18
	Unified-IO-2 XL (Lu et al., 2022)	3B	24.3	17	23.2	16	27.8	9	22.5	19	25.3	18	31.5	4	34.8	5	26.3	15
	Unified-IO-2 XXL (Lu et al., 2022)	7B	26.3	9	22.7	18	26.4	14	32.5	6	26.8	10	24.5	16	33.8	8	27.2	9
	OneLLM (Han et al., 2024)	7B	25.0	15	25.5	9	21.5	22	37.5	5	29.3	1	25.5	13	38.4	3	27.4	8
	PandaGPT (Su et al., 2023)	7B	23.5	19	23.2	16	27.6	12	45.0	1	23.8	21	28.0	12	23.9	20	26.7	13
	Video-Llama (Zhang et al., 2023a)	7B	25.5	10	22.3	19	24.4	21	30.0	9	26.2	15	25.0	14	30.7	11	26.1	17
	VideoLLaMA2 (Cheng et al., 2024)	7B	24.1	18	25.5	9	26.4	16	30.0	9	27.2	9	33.0	2	34.5	6	26.8	12
	AnyGPT (Zhan et al., 2024)	7B	24.6	16	25.0	11	26.4	17	27.5	14	29.2	2	29.0	8	25.7	17	26.1	19
	NExT-GPT (Wu et al., 2023)	7B	23.2	22	20.9	20	27.8	11	30.0	9	28.8	3	28.5	10	23.6	21	25.5	20
	VITA (Fu et al., 2024)	8 × 7B	24.1	19	26.4	8	27.8	9	22.5	19	26.3	14	31.0	6	36.8	4	26.4	14
	Qwen-2.5-Omni-7B (Xu et al., 2025)	7B	38.6	4	30.0	6	30.4	7	40.0	3	25.8	16	31.5	4	39.6	2	34.5	4
	AV-Reasoner (Lu et al., 2025)	7B	42.5	2	31.4	4	28.3	9	40.0	3	26.5	12	32.5	3	43.4	1	36.4	2
Closed Source	Gemini 1.5 Flash (Team et al., 2024a)	-	27.2	7	25.0	11	28.8	8	30.0	9	25.3	18	28.5	10	31.2	10	27.8	7
	Gemini 1.5 Flash-8B (Team et al., 2024a)	-	25.1	13	24.5	13	28.9	6	27.5	14	27.5	5	29.0	5	30.2	13	26.8	13
	Gemini 1.5 Pro (Team et al., 2024a)	-	30.8	6	31.4	4	31.3	5	37.5	5	27.7	4	20.5	21	33.0	9	30.8	6
	Gemini 2.5 Pro (Google, 2025)	-	53.6	1	40.0	1	41.9	1	32.5	6	23.8	21	40.5	1	37.8	3	44.4	1
	Reka Core (Team et al., 2024b)	67B	26.7	8	27.7	7	26.4	15	22.5	19	26.5	12	24.0	17	34.3	7	26.9	10
	Reka Flash (Team et al., 2024b)	21B	25.5	11	24.1	14	27.2	13	30.0	9	27.5	5	31.5	4	24.1	19	26.3	16
	Reka Edge (Team et al., 2024b)	7B	23.8	21	20.5	22	26.3	18	22.5	19	25.5	17	22.5	19	36.8	4	25.0	22
	GPT-4o visual caption (Hurst et al., 2024)	-	37.4	5	28.6	6	32.3	4	27.5	14	25.5	17	23.0	18	28.9	14	32.3	5
	GPT-4o audio caption (Hurst et al., 2024)	-	38.6	3	31.8	2	33.6	2	42.5	2	27.5	5	25.0	15	26.1	16	34.5	3

2023a), OneLLM (Han et al., 2024), AnyGPT (Zhan et al., 2024), NExT-GPT (Wu et al., 2023), VITA (Fu et al., 2024)). Open-source models were tested using their latest checkpoints and **default hyperparameters** from published code; closed-source models relied on official APIs.

GPT-4o Workaround. Due to API limitations preventing simultaneous multimodal inputs, we evaluate GPT-4o via two pipelines: 1. *Audio Caption Method*: We use GPT-4o-audio to generate audio captions, then use the audio caption, question text, and visual as the inputs of GPT-4o. 2. *Visual Caption Method*: We use GPT-4o to generate visual captions, then use the visual caption, question text, and audio as the inputs of GPT-4o-audio.

Baseline and Interpretation. A 25% random baseline (corresponding to a four-choice task) is established. Performance below this threshold indicates the model’s inability to tackle the task.

4.2 EXPERIMENTAL ANALYSIS

In this section, we analyze the performance of MLLMs in our AV-Odyssey benchmark, as presented in Table 4. We showcase the mean accuracy of each audio attribute. Detailed results and data distribution are provided in the Appendix. Our key findings are as follows:

1. Challenging Nature of AV-Odyssey. Table 4 shows that most MLLMs achieve barely above 25% accuracy—only marginally better than random guessing. The top-performing model, Gemini 2.5 Pro, reaches just 44.4%. These results highlight the rigor of AV-Odyssey, which evaluates capabilities beyond the scope of current training data. By establishing demanding standards, AV-Odyssey serves as a crucial benchmark for assessing MLLM robustness and versatility in audio-visual reasoning, revealing key limitations and providing insights for future advancements.

2. Comparison Between Audio Captions and Visual Captions. It can be observed that GPT-4o audio caption achieves higher performance than GPT-4o visual caption as shown in Table 4. However, drawing conclusions from this comparison is challenging, as the captions are generated by the model itself, introducing potential biases. To isolate the influence of each modality, we further investigate this phenomenon by providing ground-truth (GT) captions for both audio and visual tasks. In this setup, GPT-4o’s performance improves to 73.33% with ground-truth audio

	w/ GT audio caption	w/ GT visual caption
GPT-4o	73.33%	47.33%

Table 5: GPT-4o with GT audio/visual caption on randomly selected 300 Questions.

captions, while performance with ground-truth visual captions increases to 47.33%. These results suggest that the primary limitation of current models lies in their audio understanding capabilities.

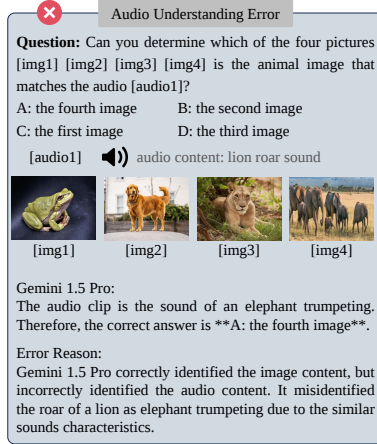


Figure 6: An example of audio understanding error. More examples are provided in the Appendix.

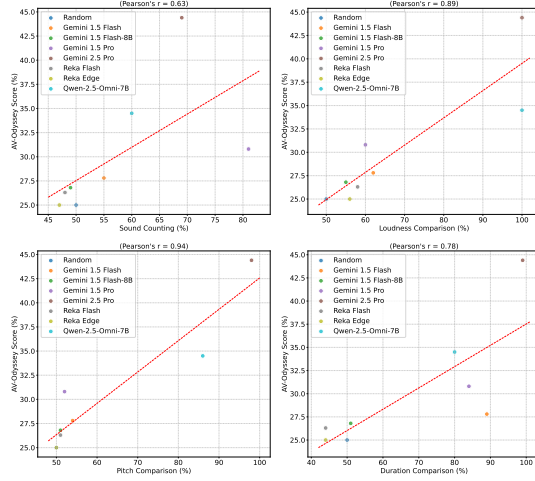


Figure 7: Comparison performance between each DeafTest task and AV-Odyssey.

3. Correlation between each DeafTest task and AV-Odyssey. We analyze the relationship between performance on DeafTest and AV-Odyssey (Fig. 7). Notably, loudness and pitch perception show a stronger correlation with AV-Odyssey performance, as these fundamental auditory components are crucial for accurate audio understanding. Tasks such as audio length recognition and sound counting exhibit weaker correlations with AV-Odyssey performance. This is because AV-Odyssey is designed to evaluate comprehensive audio-visual understanding. As such, tasks focused on isolated auditory features, such as temporal duration or discrete event counting, are less strongly aligned with the broader, more integrated tasks within AV-Odyssey.

4.3 ERROR ANALYSIS

We analyzed 104 human-annotated errors (4 randomly sampled per task) to identify failure modes. The error distribution is shown in Fig. 8, with case details in the Appendix.

1. Perception Understanding Errors (81%). Most errors stemmed from flawed input interpretation, **dominated by audio-related failures (63%)**. For example, Fig. 6 shows a misidentified audio clip leading to incorrect answers. Vision (10%) and text (8%) understanding errors were less frequent. This **aligns with DeafTest findings**: deficient audio perception undermines cross-modal integration. **2. Reasoning Errors (13%).** In these cases, Gemini 1.5 Pro correctly parsed audio/visual inputs but failed in logical inference (*e.g.*, misconnecting cause-effect relationships). **3. Other Errors (6%).** Mostly rejected responses (*e.g.*, content flagged for security).

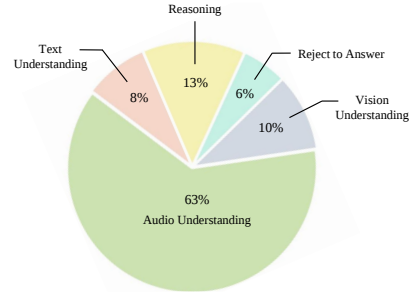


Figure 8: Distribution of 104 human-annotated errors in the Gemini 1.5 Pro.

5 LIMITATIONS AND FUTURE WORK

In this work, we introduce DeafTest and AV-Odyssey to evaluate MLLMs from fundamental audio perception and audio-visual reasoning perspectives. Further, we reveal the strong correlation between basic audio perception and complex audio-visual reasoning.

Our study does not yet address long video understanding (*e.g.*, five minutes or more), real-time interactive scenarios, or joint understanding-generation tasks. These directions remain promising avenues for future work. While our results underscore the importance of fundamental audio perception, additional factors influencing audio-visual reasoning are likely to exist and warrant further investigation.

REFERENCES

- Erhan AKBAL, Turker TUNCER, and Sengul. Vehicle interior sound dataset, October 2021. URL <https://doi.org/10.5281/zenodo.5606504>.
- Andrada. Gtzan dataset - music genre classification. <https://www.kaggle.com/datasets/andradaolteanu/gtzan-dataset-music-genre-classification>.
- Emrah AYDEMİR. Gunshot audio dataset. <https://www.kaggle.com/datasets/emrahaydemr/gunshot-audio-dataset>.
- Jinze Bai, Shuai Bai, Shusheng Yang, Shijie Wang, Sinan Tan, Peng Wang, Junyang Lin, Chang Zhou, and Jingren Zhou. Qwen-vl: A frontier large vision-language model with versatile abilities. *arXiv preprint arXiv:2308.12966*, 2023.
- Xiao Bi, Deli Chen, Guanting Chen, Shanhuang Chen, Damai Dai, Chengqi Deng, Honghui Ding, Kai Dong, Qiusi Du, Zhe Fu, et al. Deepseek llm: Scaling open-source language models with longtermism. *arXiv preprint arXiv:2401.02954*, 2024.
- Tom B Brown. Language models are few-shot learners. *arXiv preprint arXiv:2005.14165*, 2020.
- Jie Cao and Jing Xiao. An augmented benchmark dataset for geometric question answering through dual parallel text encoding. In *Proceedings of the 29th International Conference on Computational Linguistics*, pp. 1511–1520, 2022.
- Jiaqi Chen, Jianheng Tang, Jinghui Qin, Xiaodan Liang, Lingbo Liu, Eric P Xing, and Liang Lin. Geoqa: A geometric question answering benchmark towards multimodal numerical reasoning. *arXiv preprint arXiv:2105.14517*, 2021.
- Jiaqi Chen, Tong Li, Jinghui Qin, Pan Lu, Liang Lin, Chongyu Chen, and Xiaodan Liang. Unigeo: Unifying geometry logical reasoning via reformulating mathematical expression. *arXiv preprint arXiv:2212.02746*, 2022.
- Jun Chen, Deyao Zhu, Xiaoqian Shen, Xiang Li, Zechun Liu, Pengchuan Zhang, Raghuraman Krishnamoorthi, Vikas Chandra, Yunyang Xiong, and Mohamed Elhoseiny. Minigpt-v2: large language model as a unified interface for vision-language multi-task learning. *arXiv preprint arXiv:2310.09478*, 2023a.
- Keqin Chen, Zhao Zhang, Weili Zeng, Richong Zhang, Feng Zhu, and Rui Zhao. Shikra: Unleashing multimodal llm’s referential dialogue magic. *arXiv preprint arXiv:2306.15195*, 2023b.
- Lin Chen, Jinsong Li, Xiaoyi Dong, Pan Zhang, Yuhang Zang, Zehui Chen, Haodong Duan, Jiaqi Wang, Yu Qiao, Dahua Lin, et al. Are we on the right way for evaluating large vision-language models? *arXiv preprint arXiv:2403.20330*, 2024a.
- Zhe Chen, Weiyun Wang, Hao Tian, Shenglong Ye, Zhangwei Gao, Erfei Cui, Wenwen Tong, Kongzhi Hu, Jiapeng Luo, Zheng Ma, et al. How far are we to gpt-4v? closing the gap to commercial multimodal models with open-source suites. *arXiv preprint arXiv:2404.16821*, 2024b.
- Zesen Cheng, Sicong Leng, Hang Zhang, Yifei Xin, Xin Li, Guanzheng Chen, Yongxin Zhu, Wenqi Zhang, Ziyang Luo, Deli Zhao, et al. Videollama 2: Advancing spatial-temporal modeling and audio understanding in video-llms. *arXiv preprint arXiv:2406.07476*, 2024.
- Sanjoy Chowdhury, Sayan Nag, Subhrajyoti Dasgupta, Jun Chen, Mohamed Elhoseiny, Ruohan Gao, and Dinesh Manocha. Meerkat: Audio-visual large language model for grounding in space and time. In *European Conference on Computer Vision*, pp. 52–70. Springer, 2024.
- Yunfei Chu, Jin Xu, Xiaohuan Zhou, Qian Yang, Shiliang Zhang, Zhijie Yan, Chang Zhou, and Jingren Zhou. Qwen-audio: Advancing universal audio understanding via unified large-scale audio-language models. *arXiv preprint arXiv:2311.07919*, 2023.
- Yunfei Chu, Jin Xu, Qian Yang, Haojie Wei, Xipin Wei, Zhifang Guo, Yichong Leng, Yuanjun Lv, Jinzheng He, Junyang Lin, et al. Qwen2-audio technical report. *arXiv preprint arXiv:2407.10759*, 2024.

- Wenliang Dai, Junnan Li, Dongxu Li, Anthony Meng Huat Tiong, Junqi Zhao, Weisheng Wang, Boyang Li, Pascale Fung, and Steven Hoi. Instructblip: Towards general-purpose vision-language models with instruction tuning. *arXiv preprint arXiv:2305.06500*, 2023.
- Dima Damen, Hazel Doughty, Giovanni Maria Farinella, Sanja Fidler, Antonino Furnari, Evangelos Kazakos, Davide Moltisanti, Jonathan Munro, Toby Perrett, Will Price, et al. Scaling egocentric vision: The epic-kitchens dataset. In *Proceedings of the European conference on computer vision (ECCV)*, pp. 720–736, 2018.
- Michaël Defferrard, Kirell Benzi, Pierre Vandergheynst, and Xavier Bresson. Fma: A dataset for music analysis. *arXiv preprint arXiv:1612.01840*, 2016.
- Chaoyou Fu, Peixian Chen, Yunhang Shen, Yulei Qin, Mengdan Zhang, Xu Lin, Zhenyu Qiu, Wei Lin, Jinrui Yang, Xiawu Zheng, Ke Li, Xing Sun, and Rongrong Ji. Mme: A comprehensive evaluation benchmark for multimodal large language models. *arXiv preprint arXiv:2306.13394*, 2023.
- Chaoyou Fu, Haojia Lin, Zuwei Long, Yunhang Shen, Meng Zhao, Yifan Zhang, Xiong Wang, Di Yin, Long Ma, Xiawu Zheng, et al. Vita: Towards open-source interactive omni multimodal llm. *arXiv preprint arXiv:2408.05211*, 2024.
- Google. Gemini 2.5 Pro: Advanced Generative AI Model. <https://cloud.google.com/vertex-ai/generative-ai/docs/models/gemini/2-5-pro>, 2025.
- Jiaming Han, Kaixiong Gong, Yiyuan Zhang, Jiaqi Wang, Kaipeng Zhang, Dahua Lin, Yu Qiao, Peng Gao, and Xiangyu Yue. Onellm: One framework to align all modalities with language. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 26584–26595, 2024.
- Toni Heittola, Annamaria Mesaros, and Tuomas Virtanen. TAU Urban Acoustic Scenes 2019, Development dataset, March 2019. URL <https://doi.org/10.5281/zenodo.2589280>.
- Anwen Hu, Yaya Shi, Haiyang Xu, Jiabo Ye, Qinghao Ye, Ming Yan, Chenliang Li, Qi Qian, Ji Zhang, and Fei Huang. mplug-paperowl: Scientific diagram analysis with the multimodal large language model. In *Proceedings of the 32nd ACM International Conference on Multimedia*, pp. 6929–6938, 2024.
- EH Huizing. The early descriptions of the so-called tuning-fork tests of weber, rinne, schwabach, and bing: Iii. the development of the schwabach and bing tests. *ORL*, 37(2):92–96, 1975.
- Aaron Hurst, Adam Lerer, Adam P Goucher, Adam Perelman, Aditya Ramesh, Aidan Clark, AJ Ostrow, Akila Welihinda, Alan Hayes, Alec Radford, et al. Gpt-4o system card. *arXiv preprint arXiv:2410.21276*, 2024.
- Will Kay, Joao Carreira, Karen Simonyan, Brian Zhang, Chloe Hillier, Sudheendra Vijayanarasimhan, Fabio Viola, Tim Green, Trevor Back, Paul Natsev, et al. The kinetics human action video dataset. *arXiv preprint arXiv:1705.06950*, 2017.
- Sicong Leng, Yun Xing, Zesen Cheng, Yang Zhou, Hang Zhang, Xin Li, Deli Zhao, Shijian Lu, Chunyan Miao, and Lidong Bing. The curse of multi-modalities: Evaluating hallucinations of large multimodal models across language, visual, and audio. *arXiv preprint arXiv:2410.12787*, 2024.
- Bohao Li, Yuying Ge, Yi Chen, Yixiao Ge, Ruimao Zhang, and Ying Shan. Seed-bench-2-plus: Benchmarking multimodal large language models with text-rich visual comprehension. *arXiv preprint arXiv:2404.16790*, 2024a.
- Bohao Li, Yuying Ge, Yixiao Ge, Guangzhi Wang, Rui Wang, Ruimao Zhang, and Ying Shan. Seed-bench: Benchmarking multimodal large language models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 13299–13308, 2024b.
- Guangyao Li, Yake Wei, Yapeng Tian, Chenliang Xu, Ji-Rong Wen, and Di Hu. Learning to answer questions in dynamic audio-visual scenarios. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 19108–19118, 2022.

- Junnan Li, Dongxu Li, Silvio Savarese, and Steven Hoi. Blip-2: Bootstrapping language-image pre-training with frozen image encoders and large language models. *ICML*, 2023a.
- Yanwei Li, Chengyao Wang, and Jiaya Jia. Llama-vid: An image is worth 2 tokens in large language models. In *European Conference on Computer Vision*, 2024c.
- Yifan Li, Yifan Du, Kun Zhou, Jinpeng Wang, Wayne Xin Zhao, and Ji-Rong Wen. Evaluating object hallucination in large vision-language models. *arXiv preprint arXiv:2305.10355*, 2023b.
- Yizhi Li, Ge Zhang, Yinghao Ma, Ruibin Yuan, Kang Zhu, Hangyu Guo, Yiming Liang, Jiaheng Liu, Jian Yang, Siwei Wu, et al. Omnibench: Towards the future of universal omni-language models. *NeurIPS*, 2025.
- Ji Lin, Hongxu Yin, Wei Ping, Pavlo Molchanov, Mohammad Shoeybi, and Song Han. Vila: On pre-training for visual language models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 26689–26699, 2024.
- Haotian Liu, Chunyuan Li, Yuheng Li, and Yong Jae Lee. Improved baselines with visual instruction tuning. *arXiv preprint arXiv:2310.03744*, 2023a.
- Haotian Liu, Chunyuan Li, Qingyang Wu, and Yong Jae Lee. Visual instruction tuning. *arXiv preprint arXiv:2304.08485*, 2023b.
- Jiajun Liu, Yibing Wang, Hanghang Ma, Xiaoping Wu, Xiaoqi Ma, xiaoming Wei, Jianbin Jiao, Enhua Wu, and Jie Hu. Kangaroo: A powerful video-language model supporting long-context video input. *arXiv preprint arXiv:2408.15542*, 2024.
- Yuan Liu, Haodong Duan, Yuanhan Zhang, Bo Li, Songyang Zhang, Wangbo Zhao, Yike Yuan, Jiaqi Wang, Conghui He, Ziwei Liu, et al. Mmbench: Is your multi-modal model an all-around player? *arXiv preprint arXiv:2307.06281*, 2023c.
- Eu Jin Lok. Toronto emotional speech set (tess). <https://www.kaggle.com/datasets/ejlok1/toronto-emotional-speech-set-tess>.
- Jiasen Lu, Christopher Clark, Rowan Zellers, Roozbeh Mottaghi, and Aniruddha Kembhavi. Unified-io: A unified model for vision, language, and multi-modal tasks. In *The Eleventh International Conference on Learning Representations*, 2022.
- Lidong Lu, Guo Chen, Zhiqi Li, Yicheng Liu, and Tong Lu. Av-reasoner: Improving and benchmarking clue-grounded audio-visual counting for mllms. 2025.
- Pan Lu, Hritik Bansal, Tony Xia, Jiacheng Liu, Chunyuan Li, Hannaneh Hajishirzi, Hao Cheng, Kai-Wei Chang, Michel Galley, and Jianfeng Gao. Mathvista: Evaluating mathematical reasoning of foundation models in visual contexts. *arXiv preprint arXiv:2310.02255*, 2023.
- Tengchao Lv, Yupan Huang, Jingye Chen, Lei Cui, Shuming Ma, Yaoyao Chang, Shaohan Huang, Wenhui Wang, Li Dong, Weiyao Luo, et al. Kosmos-2.5: A multimodal literate model. *arXiv preprint arXiv:2309.11419*, 2023.
- Jeannette Shijie Ma. Eating sound collection. <https://www.kaggle.com/datasets/mashijie/eating-sound-collection>.
- Potsawee Manakul, Guangzhi Sun, Warit Sirichotedumrong, Kasima Tharnpipitchai, and Kunat Pipatanakul. Enhancing low-resource language and instruction following capabilities of audio language models. *arXiv preprint arXiv:2409.10999*, 2024.
- Microsoft. Birdclef 2020. <https://www.imageclef.org/BirdCLEF2020>.
- MISC. Music instrument sounds for classification. <https://www.kaggle.com/datasets/abdulvahap/music-instrument-sounds-for-classification>.
- OpenAI. Hello gpt-4o. <https://openai.com/index/hello-gpt-4o/>.
- R OpenAI. Gpt-4 technical report. arxiv 2303.08774. *View in Article*, 2(5), 2023.

- Fabian Ostermann, Igor Vatolkin, and Martin Ebeling. Aam: a dataset of artificial audio multitracks for diverse music information retrieval tasks. *EURASIP Journal on Audio, Speech, and Music Processing*, 2023(1):13, 2023.
- Zhiliang Peng, Wenhui Wang, Li Dong, Yaru Hao, Shaohan Huang, Shuming Ma, and Furu Wei. Kosmos-2: Grounding multimodal large language models to the world. *arXiv preprint arXiv:2306.14824*, 2023.
- Rushi Balaji Putthewad. Sound classification of animal voice. <https://www.kaggle.com/datasets/rushibalajiputthewad/sound-classification-of-animal-voice>.
- Alec Radford. Improving language understanding by generative pre-training. 2018.
- Pooyan Rahmanzadehgervi, Logan Bolton, Mohammad Reza Taesiri, and Anh Totti Nguyen. Vision language models are blind. *arXiv preprint arXiv:2407.06581*, 2024.
- Shuhuai Ren, Linli Yao, Shicheng Li, Xu Sun, and Lu Hou. Timechat: A time-sensitive multimodal large language model for long video understanding. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 14313–14323, 2024.
- Minjoon Seo, Hannaneh Hajishirzi, Ali Farhadi, Oren Etzioni, and Clint Malcolm. Solving geometry problems: Combining text and diagram interpretation. In *Proceedings of the 2015 conference on empirical methods in natural language processing*, pp. 1466–1476, 2015.
- Xiaoqian Shen, Yuniang Xiong, Changsheng Zhao, Lemeng Wu, Jun Chen, Chenchen Zhu, Zechun Liu, Fanyi Xiao, Balakrishnan Varadarajan, Florian Bordes, Zhuang Liu, Hu Xu, Hyunwoo J. Kim, Bilge Soran, Raghuraman Krishnamoorthi, Mohamed Elhoseiny, and Vikas Chandra. Longvu: Spatiotemporal adaptive compression for long video-language understanding. *arXiv preprint arXiv:2410.17434*, 2024.
- Kazuki Shimada, Archontis Politis, Parthasaarathy Sudarsanam, Daniel A Krause, Kengo Uchida, Sharath Adavanne, Aapo Hakala, Yuichiro Koyama, Naoya Takahashi, Shusuke Takahashi, et al. Starss23: An audio-visual dataset of spatial recordings of real scenes with spatiotemporal annotations of sound events. *Advances in Neural Information Processing Systems*, 36, 2024.
- K Soomro. Ucf101: A dataset of 101 human actions classes from videos in the wild. *arXiv preprint arXiv:1212.0402*, 2012.
- Auston Sterling, Justin Wilson, Sam Lowe, and Ming C Lin. Isnn: Impact sound neural network for audio-visual object classification. In *Proceedings of the European Conference on Computer Vision (ECCV)*, pp. 555–572, 2018.
- Yixuan Su, Tian Lan, Huayang Li, Jialu Xu, Yan Wang, and Deng Cai. Pandagpt: One model to instruction-follow them all. *arXiv preprint arXiv:2305.16355*, 2023.
- Guangzhi Sun, Wenyi Yu, Changli Tang, Xianzhao Chen, Tian Tan, Wei Li, Lu Lu, Zejun Ma, Yuxuan Wang, and Chao Zhang. video-salmonn: Speech-enhanced audio-visual large language models. *arXiv preprint arXiv:2406.15704*, 2024.
- Kim Sung-Bin, Oh Hyun-Bin, Jungmok Lee, Arda Senocak, Joon Son Chung, and Tae-Hyun Oh. Avhbench: A cross-modal hallucination benchmark for audio-visual large language models. *arXiv preprint arXiv:2410.18325*, 2024.
- Changli Tang, Wenyi Yu, Guangzhi Sun, Xianzhao Chen, Tian Tan, Wei Li, Lu Lu, Zejun Ma, and Chao Zhang. Salmonn: Towards generic hearing abilities for large language models. *arXiv preprint arXiv:2310.13289*, 2023.
- Yunlong Tang, Daiki Shimada, Jing Bi, and Chenliang Xu. Avicuna: Audio-visual llm with interleaver and context-boundary alignment for temporal referential dialogue. *arXiv e-prints*, pp. arXiv-2403, 2024.

- Gemini Team, Petko Georgiev, Ving Ian Lei, Ryan Burnell, Libin Bai, Anmol Gulati, Garrett Tanzer, Damien Vincent, Zhufeng Pan, Shibo Wang, et al. Gemini 1.5: Unlocking multimodal understanding across millions of tokens of context. *arXiv preprint arXiv:2403.05530*, 2024a.
- Reka Team, Aitor Ormazabal, Che Zheng, Cyprien de Masson d’Autume, Dani Yogatama, Deyu Fu, Donovan Ong, Eric Chen, Eugénie Lamprecht, Hai Pham, et al. Reka core, flash, and edge: A series of powerful multimodal language models. *arXiv preprint arXiv:2404.12387*, 2024b.
- Yapeng Tian, Jing Shi, Bochen Li, Zhiyao Duan, and Chenliang Xu. Audio-visual event localization in unconstrained videos. In *Proceedings of the European conference on computer vision (ECCV)*, pp. 247–263, 2018.
- Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, et al. Llama: Open and efficient foundation language models. *arXiv preprint arXiv:2302.13971*, 2023.
- Shuhei Tsuchida, Satoru Fukayama, Masahiro Hamasaki, and Masataka Goto. Aist dance video database: Multi-genre, multi-dancer, and multi-camera database for dance information processing. In *Proceedings of the 20th International Society for Music Information Retrieval Conference, ISMIR 2019*, Delft, Netherlands, November 2019.
- Peng Wang, Shuai Bai, Sinan Tan, Shijie Wang, Zhihao Fan, Jinze Bai, Keqin Chen, Xuejing Liu, Jialin Wang, Wenbin Ge, et al. Qwen2-vl: Enhancing vision-language model’s perception of the world at any resolution. *arXiv preprint arXiv:2409.12191*, 2024.
- Weihan Wang, Qingsong Lv, Wenmeng Yu, Wenyi Hong, Ji Qi, Yan Wang, Junhui Ji, Zhuoyi Yang, Lei Zhao, Xixuan Song, et al. Cogvlm: Visual expert for pretrained language models. *arXiv preprint arXiv:2311.03079*, 2023.
- Shengqiong Wu, Hao Fei, Leigang Qu, Wei Ji, and Tat-Seng Chua. Next-gpt: Any-to-any multimodal llm. *arXiv preprint arXiv:2309.05519*, 2023.
- Jin Xu, Zhifang Guo, Jinzheng He, Hangrui Hu, Ting He, Shuai Bai, Keqin Chen, Jialin Wang, Yang Fan, and Kai Dang. Qwen2.5-omni technical report. 2025.
- Jingyi Xu, Hieu Le, Vu Nguyen, Viresh Ranjan, and Dimitris Samaras. Zero-shot object counting. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 15548–15557, 2023.
- Le Xue, Manli Shu, Anas Awadalla, Jun Wang, An Yan, Senthil Purushwalkam, Honglu Zhou, Viraj Prabhu, Yutong Dai, Michael S Ryoo, et al. xgen-mm (blip-3): A family of open large multimodal models. *arXiv preprint arXiv:2408.08872*, 2024.
- Pinci Yang, Xin Wang, Xuguang Duan, Hong Chen, Runze Hou, Cong Jin, and Wenwu Zhu. Avqa: A dataset for audio-visual question answering on videos. In *Proceedings of the 30th ACM international conference on multimedia*, pp. 3480–3491, 2022.
- Zhengyuan Yang, Linjie Li, Kevin Lin, Jianfeng Wang, Chung-Ching Lin, Zicheng Liu, and Lijuan Wang. The dawn of Imms: Preliminary explorations with gpt-4v (ision). *arXiv preprint arXiv:2309.17421*, 9(1):1, 2023.
- Yuan Yao, Tianyu Yu, Ao Zhang, Chongyi Wang, Junbo Cui, Hongji Zhu, Tianchi Cai, Haoyu Li, Weilin Zhao, Zhihui He, et al. Minicpm-v: A gpt-4v level mllm on your phone. *arXiv preprint arXiv:2408.01800*, 2024.
- yash bhaskar. Emotify - emotion classification in songs. <https://www.kaggle.com/datasets/yash9439/emotify-emotion-classification-in-songs>.
- Qinghao Ye, Haiyang Xu, Guohai Xu, Jiabo Ye, Ming Yan, Yiyang Zhou, Junyang Wang, Anwen Hu, Pengcheng Shi, Yaya Shi, et al. mplug-owl: Modularization empowers large language models with multimodality. *arXiv preprint arXiv:2304.14178*, 2023.

- Xiang Yue, Yuansheng Ni, Kai Zhang, Tianyu Zheng, Ruoqi Liu, Ge Zhang, Samuel Stevens, Dongfu Jiang, Weiming Ren, Yuxuan Sun, et al. Mmmu: A massive multi-discipline multi-modal understanding and reasoning benchmark for expert agi. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 9556–9567, 2024.
- Heeseung Yun, Youngjae Yu, Wonsuk Yang, Kangil Lee, and Gunhee Kim. Pano-avqa: Grounded audio-visual question answering on 360deg videos. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pp. 2031–2041, 2021.
- Jun Zhan, Junqi Dai, Jiasheng Ye, Yunhua Zhou, Dong Zhang, Zhigeng Liu, Xin Zhang, Ruibin Yuan, Ge Zhang, Linyang Li, et al. Anygpt: Unified multimodal llm with discrete sequence modeling. *arXiv preprint arXiv:2402.12226*, 2024.
- Hang Zhang, Xin Li, and Lidong Bing. Video-llama: An instruction-tuned audio-visual language model for video understanding. *arXiv preprint arXiv:2306.02858*, 2023a.
- Pan Zhang, Xiaoyi Dong Bin Wang, Yuhang Cao, Chao Xu, Linke Ouyang, Zhiyuan Zhao, Shuangrui Ding, Songyang Zhang, Haodong Duan, Hang Yan, et al. Internlm-xcomposer: A vision-language large model for advanced text-image comprehension and composition. *arXiv preprint arXiv:2309.15112*, 2023b.
- Peiyuan Zhang, Kaichen Zhang, Bo Li, Guangtao Zeng, Jingkang Yang, Yuanhan Zhang, Ziyue Wang, Haoran Tan, Chunyuan Li, and Ziwei Liu. Long context transfer from language to vision. *arXiv preprint arXiv:2406.16852*, 2024a. URL <https://arxiv.org/abs/2406.16852>.
- Renrui Zhang, Dongzhi Jiang, Yichi Zhang, Haokun Lin, Ziyu Guo, Pengshuo Qiu, Aojun Zhou, Pan Lu, Kai-Wei Chang, Yu Qiao, et al. Mathverse: Does your multi-modal llm truly see the diagrams in visual math problems? In *European Conference on Computer Vision*, pp. 169–186. Springer, 2025.
- Yanzhe Zhang, Ruiyi Zhang, Jiuxiang Gu, Yufan Zhou, Nedim Lipka, Diyi Yang, and Tong Sun. Lllavar: Enhanced visual instruction tuning for text-rich image understanding, 2023c.
- Yu Zhang, Changhao Pan, Wenxiang Guo, Ruiqi Li, Zhiyuan Zhu, Jiale Wang, Wenhao Xu, Jingyu Lu, Zhiqing Hong, Chuxin Wang, et al. Gtsinger: A global multi-technique singing corpus with realistic music scores for all singing tasks. *arXiv preprint arXiv:2409.13832*, 2024b.
- Deyao Zhu, Jun Chen, Xiaoqian Shen, Xiang Li, and Mohamed Elhoseiny. Minigpt-4: Enhancing vision-language understanding with advanced large language models. *arXiv preprint arXiv:2304.10592*, 2023.

A APPENDIX

B USE OF LLM

We employed large language models (LLMs), such as ChatGPT, to assist in manuscript preparation. Their role was restricted to language refinement (including grammar, spelling, and word choice), code formatting (e.g., adding explanatory comments), and drafting preliminary figures to support the development of final visualizations. All scientific concepts, analyses, and conclusions were independently conceived, verified, and interpreted by the authors. We sincerely acknowledge the valuable support of LLMs in this work.

C DATA DISTRIBUTION

In this section, we present the detailed data distribution of our AV-Odyssey Bench in Table 6. Our AV-Odyssey bench consists of 26 tasks covering a wide range of task categories.

We will make all the data and evaluation codes public.

Table 6: Detailed task statistics in AV-Odyssey Bench.

Task ID	Task Name	Task Category	Class	Number	Data Source
1	Instrument Recognition	Timbre	28	200	MISC
2	Singer Recognition	Timbre	20	200	Internet
3	Gunshot Recognition	Timbre	13	200	AYDEMİR
4	Bird Recognition	Timbre	39	200	Microsoft
5	Animal Recognition	Timbre	13	200	Putthewad
6	Transportation Recognition	Timbre	8	200	AKBAL et al. (2021)
7	Material Recognition	Timbre	10	200	Sterling et al. (2018)
8	Scene Recognition	Timbre	8	200	Heittola et al. (2019)
9	Hazard Recognition	Timbre	8	108	Kay et al. (2017)
10	Action Recognition	Timbre	20	196	Soomro (2012)
11	Eating Sound Recognition	Timbre	20	200	Ma
12	Speech Sentiment Analysis	Tone	7	200	Lok
13	Meme Understanding	Tone	N/A	20	Internet
14	Music Sentiment Analysis	Melody	7	197	yash bhaskar
15	Music Genre Classification	Melody	8	200	Andrada
16	Dance and Music Matching	Melody	10	200	Tsuchida et al. (2019)
17	Film and Music Matching	Melody	5	200	Defferrard et al. (2016)
18	Music Score Matching	Melody	N/A	200	Zhang et al. (2024b)
19	Audio 3D Angle Estimation	Space	N/A	20	Shimada et al. (2024)
20	Audio Distance Estimation	Space	N/A	20	Shimada et al. (2024)
21	Audio Time Estimation	Time	N/A	200	Damen et al. (2018)
22	Audio-Visual Synchronization	Time	N/A	200	Tian et al. (2018)
23	Action Sequencing	Time	N/A	200	Damen et al. (2018)
24	Hallucination Evaluation	Hallucination	19	200	Ostermann et al. (2023)
25	Action Prediction	Intricacy	N/A	199	Damen et al. (2018)
26	Action Tracing	Intricacy	N/A	195	Damen et al. (2018)

D BREAKDOWN RESULTS

In this section, we provide detailed results of evaluated methods on our proposed AV-Odyssey Bench, as demonstrated in Table 7 and Table 8.

Table 7: Evaluation results of various MLLMs in ‘Timbre’ part of AV-Odyssey Bench. The best performance is in bold. The corresponding brackets for each task indicate the number of associated questions.

	Model	LLM Size	Instrument Recognition	Singer Recognition	Gunshot Recognition	Bird Recognition	Animal Recognition	Transportation Recognition	Material Recognition	Scene Recognition	Hazard Recognition	Action Recognition	Eating Sound Recognition
			(200)	(200)	(200)	(200)	(200)	(200)	(200)	(200)	(108)	(196)	(200)
Open Source	Unified-IO-2 L (Lu et al., 2022)	1B	20.5	22.5	25.5	18.5	27.0	26.5	23.0	28.0	21.3	20.9	26.5
	Unified-IO-2 XL (Lu et al., 2022)	3B	20.0	23.5	24.0	20.5	27.5	26.0	27.5	30.0	19.4	19.9	26.5
	Unified-IO-2 XXL (Lu et al., 2022)	7B	29.5	24.0	23.5	29.0	23.5	25.5	30.5	26.5	23.1	27.0	25.5
	OneLLM (Han et al., 2024)	7B	26.0	21.5	27.0	26.0	22.0	20.0	29.5	24.5	26.9	23.0	29.5
	PandaGPT (Su et al., 2023)	7B	20.0	21.5	23.0	17.5	26.0	26.5	28.0	27.0	23.1	21.4	24.5
	Video-llama (Zhang et al., 2023a)	7B	22.5	24.5	27.0	26.5	27.0	23.5	28.0	25.0	25.0	26.0	25.5
	VideoLaMA2 (Cheng et al., 2024)	7B	22.5	24.0	27.0	17.0	23.5	27.5	26.5	26.5	19.4	23.0	25.5
	AnyGPT (Zhan et al., 2024)	7B	22.5	28.5	28.0	17.5	24.0	25.5	23.0	28.0	25.9	20.4	27.5
	NExT-GPT (Wu et al., 2023)	7B	21.0	23.5	25.5	21.5	25.5	25.5	21.0	24.0	19.4	23.0	24.0
	VITA (Fu et al., 2024)	8 × 7B	22.0	20.5	24.5	21.5	27.5	25.0	23.5	28.5	21.3	19.4	29.5
	Qwen-2.5-Omni-7B (Xu et al., 2025)	7B	52.0	26.0	24.0	28.5	66.0	34.0	32.5	33.5	45.4	61.7	23.0
	AV-Reasoner (Lu et al., 2025)	7B	56.5	28.5	32.0	33.5	53.5	48.5	36.5	39.5	47.2	67.9	27.0
Closed Source	Gemini 1.5 Flash (Team et al., 2024a)	-	24.5	24.0	23.5	17.0	32.5	26.0	22.5	29.5	34.3	48.0	21.5
	Gemini 1.5 Flash-8B (Team et al., 2024a)	-	16.5	22.5	24.0	19.0	28.0	26.5	27.0	29.0	26.9	32.7	24.5
	Gemini 1.5 Pro (Team et al., 2024a)	-	33.0	26.0	29.0	25.0	25.5	26.0	29.5	30.0	38.0	57.7	22.5
	Gemini 2.5 Pro (Google, 2025)	-	79.0	41.0	39.5	21.0	76.0	60.5	39.0	44.0	65.7	87.8	42.5
	Reka Core (Team et al., 2024b)	67B	20.0	26.5	25.0	24.0	27.0	30.0	27.0	27.0	25.0	34.2	21.5
	Reka Flash (Team et al., 2024b)	21B	20.0	22.5	26.5	26.0	28.5	26.5	26.5	29.0	28.7	22.4	25.0
	Reka Edge (Team et al., 2024b)	7B	21.5	24.0	30.5	20.0	19.5	22.5	20.5	25.5	25.9	23.5	29.0
	GPT-4o visual caption (Hurst et al., 2024)	-	33.0	30.5	24.0	26.5	43.0	42.0	32.5	39.0	49.1	67.3	30.5
	GPT-4o audio caption (Hurst et al., 2024)	-	40.0	38.0	27.5	26.5	45.0	42.0	27.0	41.0	42.6	62.2	35.5

Table 8: Evaluation results of various MLLMs in ‘Time’, ‘Melody’, ‘Space’. ‘Time’, ‘Hallucination’, and ‘Intracacy’ parts of AV-Odyssey Bench. The best (second best) is in bold (underline). The corresponding brackets for each task indicate the number of associated questions.

	Model	LLM Size	Tone		Melody							Space		Time		Hallucination		Intracacy	
			Speech Analysis	Sentiment Understanding	Music Analysis	Sentiment Classification	Music Genre Classification	Dance and Music Matching	Music Film Matching	Music Score Matching	Audio 3D Estimation	Angle Audio Estimation	Distance Audio Estimation	Audio Time Estimation	Audio-Visual Synchronization	Action Sequencing	Hallucination Evaluation	Action Prediction	Tracing
			(200)	(20)	(97)	(200)	(200)	(200)	(200)	(200)	(20)	(20)	(20)	(200)	(200)	(200)	(200)	(199)	(195)
Open Source	Unified-IO-2 L (Lu et al., 2022)	1B	24.5	20.0	22.9	31.0	27.5	32.5	24.5	15.0	15.0	15.0	28.0	25.5	27.0	30.0	27.1	33.8	
	Unified-IO-2 XL (Lu et al., 2022)	3B	23.0	25.0	26.9	30.5	27.0	31.5	22.5	30.0	15.0	15.0	26.5	25.5	24.0	<u>31.5</u>	35.7	33.8	
	Unified-IO-2 XXL (Lu et al., 2022)	7B	23.0	20.0	23.9	31.5	27.5	24.5	23.5	80.0	15.0	28.0	25.0	27.5	24.5	27.5	33.2	34.4	
	OneLLM (Han et al., 2024)	7B	26.0	20.0	20.8	23.5	26.5	18.5	18.0	<u>45.0</u>	30.0	31.5	29.5	27.0	25.5	41.7	34.9		
	PandaGPT (Su et al., 2023)	7B	23.5	20.0	21.6	28.0	27.0	32.5	26.0	<u>45.0</u>	45.0	18.5	26.0	27.0	28.0	19.6	28.2		
	Video-llama (Zhang et al., 2023a)	7B	23.0	15.0	25.8	24.0	20.0	25.0	28.0	<u>45.0</u>	15.0	28.5	23.5	26.5	25.0	28.6	33.8		
	VideoLaMA2 (Cheng et al., 2024)	7B	26.0	20.0	26.8	29.0	25.5	30.5	20.5	<u>45.0</u>	15.0	28.5	26.5	26.5	33.0	28.6	40.5		
	AnyGPT (Zhan et al., 2024)	7B	25.5	20.0	23.4	29.5	25.5	26.0	26.0	40.0	15.0	30.5	<u>24.0</u>	29.0	29.0	21.1	30.3		
	NExT-GPT (Wu et al., 2023)	7B	21.5	15.0	23.7	26.0	28.0	31.0	28.0	<u>45.0</u>	15.0	31.5	24.0	31.0	28.5	20.6	26.7		
	VITA (Fu et al., 2024)	8 × 7B	24.5	45.0	26.8	26.0	27.5	33.5	24.5	25.0	20.0	26.5	25.5	27.0	31.0	34.2	<u>39.5</u>		
	Qwen-2.5-Omni-7B (Xu et al., 2025)	7B	30.5	25.0	24.4	47.5	23.5	30.0	26.5	50.0	30.0	24.5	28.5	24.5	31.5	39.7	39.5		
	AV-Reasoner (Lu et al., 2025)	7B	29.0	55.0	22.8	40.5	24.5	31.5	22.5	55.0	25.0	35.0	27.5	27.0	32.5	51.3	35.4		
Closed Source	Gemini 1.5 Flash (Team et al., 2024a)	-	23.5	40.0	21.3	31.0	27.5	32.5	28.0	30.0	30.0	27.5	23.5	25.0	28.5	27.6	34.9		
	Gemini 1.5 Flash-8B (Team et al., 2024a)	-	24.5	25.0	25.9	33.0	27.5	32.0	24.5	40.0	15.0	31.0	25.5	26.0	29.0	25.6	34.9		
	Gemini 1.5 Pro (Team et al., 2024a)	-	29.5	50.0	25.4	42.5	28.0	28.5	29.0	35.0	<u>40.0</u>	30.0	24.5	28.5	20.5	32.2	33.8		
	Gemini 2.5 Pro (Google, 2025)	-	38.5	55.0	27.4	79.5	24.5	50.5	27.5	45.0	20.0	18.5	19.5	33.5	40.5	38.2	37.4		
	Reka Core (Team et al., 2024b)	67B	<u>28.5</u>	20.0	22.8	24.5	27.5	30.0	25.5	25.0	20.0	30.0	25.5	24.0	24.0	33.7	34.9		
	Reka Flash (Team et al., 2024b)	21B	24.5	20.0	30.5	29.5	27.5	25.5	24.5	<u>45.0</u>	15.0	30.0	25.5	27.0	<u>31.5</u>	19.1	29.2		
	Reka Edge (Team et al., 2024b)	7B	20.5	20.0	24.9	24.5	27.5	20.0	24.0	30.0	15.0	30.0	25.5	21.0	22.5	38.2	35.4		
	GPT-4o visual caption (Hurst et al., 2024)	-	26.0	55.0	24.4	<u>48.0</u>	27.0	34.5	23.5	25.0	30.0	21.5	22.5	<u>32.5</u>	23.0	32.2	25.6		
	GPT-4o audio caption (Hurst et al., 2024)	-	28.0	70.0	24.4	56.5	27.5	32.5	22.5	30.0	35.0	23.5	25.5	33.5	25.0	30.2	22.0		

E CASE STUDY

LIST OF FIGURES

1	Illustration of two out of four DeafTest tasks: loudness and pitch comparison. . . .	3
2	Performance on DeafTest and AV-Odyssey, showing a strong positive correlation. .	3
3	Overview of AV-Odyssey Benchmark. AV-Odyssey Bench demonstrates three major features: 1. Comprehensive Audio Attributes; 2. Extensive Domains; 3. Interleaved Text, Audio, and Images.	3
4	Overview of the 26 evaluation tasks, categorized into 7 classes based on sound attributes.	5
5	Sampled examples from our AV-Odyssey Benchmark.	6
6	An example of audio understanding error. More examples are provided in the Appendix.	9
7	Comparison performance between each DeafTest task and AV-Odyssey.	9
8	Distribution of 104 human-annotated errors in the Gemini 1.5 Pro.	9
9	A sampled error case in the instrument recognition task.	20
10	A sampled error case in the singer recognition task.	21
11	A sampled error case in the gunshot recognition task.	22
12	A sampled error case in the bird recognition task.	23
13	A sampled error case in the animal recognition task.	24
14	A sampled error case in the transportation recognition task.	25
15	A sampled error case in the material recognition task.	26
16	A sampled error case in the scene recognition task.	27
17	A sampled error case in the hazard recognition task.	28
18	A sampled error case in the action recognition task.	29
19	A sampled error case in the eating sound recognition task.	30
20	A sampled error case in the speech sentiment analysis task.	31
21	A sampled error case in the meme understanding task.	32
22	A sampled error case in the music sentiment analysis task.	33
23	A sampled error case in the music genre classification task.	34
24	A sampled error case in the dance and music matching task.	35
25	A sampled error case in the film and music matching task.	36
26	A sampled error case in the music score matching task.	37
27	A sampled error case in the audio 3D angle estimation task.	38
28	A sampled error case in the audio distance estimation task.	39
29	A sampled error case in the audio time estimation task.	40
30	A sampled error case in the audio-visual synchronization task.	41
31	A sampled error case in the action sequencing task.	42
32	A sampled error case in the hallucination evaluation task.	43
33	A sampled error case in the action prediction task.	44

972	34	A sampled error case in the action tracing task.	45
973			
974			
975			
976			
977			
978			
979			
980			
981			
982			
983			
984			
985			
986			
987			
988			
989			
990			
991			
992			
993			
994			
995			
996			
997			
998			
999			
1000			
1001			
1002			
1003			
1004			
1005			
1006			
1007			
1008			
1009			
1010			
1011			
1012			
1013			
1014			
1015			
1016			
1017			
1018			
1019			
1020			
1021			
1022			
1023			
1024			
1025			

A: the second image
B: the fourth image
C: the third image
D: the first image



Error Reason: Gemini 1.5 Pro accurately identified the visual content in the image but made an error in audio recognition, incorrectly attributing Rihanna's vocal timbre to Billie Eilish.

Figure 10: A sampled error case in the singer recognition task.

A: the second audio
B: the third audio
C: the fourth audio
D: the first audio



© Randy T. Karkner

Error Category: Audio Understanding Error, Vision Understanding Error
Error Reason: Gemini 1.5 Pro made errors in both visual and audio recognition: it incorrectly identified the common grackle bird in the image as a Brewer's Blackbird and subsequently misattributed the sound of the common grackle to that of a Brewer's Blackbird.

Figure 12: A sampled error case in the bird recognition task.

A: the fourth image
B: the second image
C: the first image
D: the third image



Correct Answer: B: the second image

25

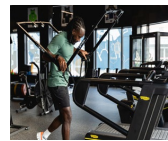
A: the third image
B: the second image
C: the fourth image
D: the first image



Error Reason: Gemini 1.5 Pro correctly identified the content of the image but made an error in audio recognition, mistakenly identifying the sound of fire burning as the sound of water boiling.

Figure 17: A sampled error case in the hazard recognition task.

A: the fourth image
B: the second image
C: the first image
D: the third image



Error Reason: Gemini 1.5 Pro correctly identified the content of the image but made an error in audio recognition, mistakenly identifying the sound of running on a treadmill as the sound of playing basketball.

Figure 18: A sampled error case in the action recognition task.

1566
1567
1568
1569
1570
1571
1572
1573
1574
1575
1576
1577
1578
1579
1580
1581
1582
1583
1584
1585
1586
1587
1588
1589
1590
1591
1592
1593
1594
1595
1596
1597
1598
1599
1600
1601
1602
1603
1604
1605
1606
1607
1608
1609
1610
1611
1612
1613
1614
1615
1616
1617
1618
1619

Timbre: Eating Sound Recognition

Question: From the foods in the images [img1], [img2], [img3], or [img4], which one is most likely to create the sound [audio1] when consumed?

A: the fourth image
B: the first image
C: the second image
D: the third image

[audio1]

audio content: the sound of eating juicy grapes

[img1]

[img2]

[img3]

[img4]

Gemini 1.5 Pro: The answer is D: the third image (French fries). The sound "crunch" is associated with crispy or brittle foods. French fries, when properly cooked, have a crisp exterior that makes a crunching sound when bitten into. The other foods are less likely to produce this specific sound. Grapes are juicy and soft, chocolate is typically smooth or melts in the mouth, and noodles, though they can have some texture, don't typically produce a distinct crunch.

Error Category: **Audio Understanding Error**

Error Reason: Gemini 1.5 Pro correctly identified the content of the image but made an error in audio recognition, mistakenly identifying the sound of eating juicy grapes as the sound of eating crispy chips.

Correct Answer: B: the first image

Figure 19: A sampled error case in the eating sound recognition task.

Tone: Speech Sentiment Analysis

Question: Which of the following images [img1] [img2] [img3] [img4] best represents the emotion conveyed in this audio clip [audio1]?

A: the fourth image

B: the second image

C: the third image

D: the first image

[audio1]  audio content: an angry voice

[img1] 

[img2] 

[img3] 

[img4] 

Gemini 1.5 Pro: The audio clip expresses disgust. The woman in the first image best depicts this emotion. So the answer is D.

Error Category: **Audio Understanding Error**
Error Reason: Gemini 1.5 Pro correctly identified the content of the image but made an error in audio recognition, mistakenly identifying an angry voice as a disgusted voice.

Correct Answer: C: the third image

Figure 20: A sampled error case in the speech sentiment analysis task.

31

A: The humor comes from the cat's overly excited face, which contradicts the serious tone of the sound, creating a funny visual and audio mismatch.

B: The meme is funny because the cat's expression is overly dramatic, and the sound is yelled, adding intensity to simple or mundane events.

C: The humor comes from the perfect match between the cat's blank, emotionless stare and the disinterested sound, both of which represent a complete lack of energy or surprise, making even absurd situations seem trivial.

D: It's funny because the cat's face shows anger, but the sound is very calm, making the reaction seem less serious than expected in a high-stress scenario.

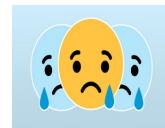


Error Category: Reject to Answer

Correct Answer: C

32

A: the fourth image
B: the second image
C: the first image
D: the third image



Error Reason: Gemini 1.5 Pro correctly identified the content of the image but made an error in audio recognition, mistakenly identifying joyful music as sad music.

Correct Answer: C: the first image

33

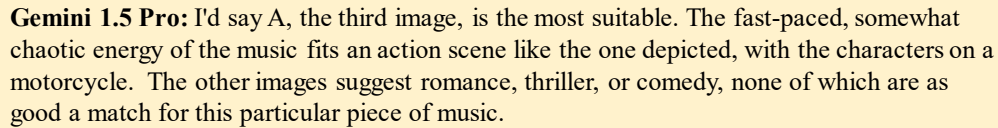
A: the fourth image
B: the second image
C: the first image
D: the third image



Error Reason: Gemini 1.5 Pro correctly identified the content of the image but made an error in audio recognition, mistakenly identifying classical music as country music.

Figure 23: A sampled error case in the music genre classification task.

[audio1]  audio content: fast-paced, cheerful music accompanied by comical screams



Error Reason: Gemini 1.5 Pro successfully understood both the audio and visual inputs but made a reasoning error when interpreting the data. It incorrectly associated the fast-paced nature of the music with the action genre, assuming that the combination of high energy and disorder aligns with action films. However, the audio, characterized by fast-paced music accompanied by comical screams, is more indicative of a comedy film rather than an action movie. This mistake highlights a gap in the model's contextual reasoning, as it overlooked the playful and humorous tone of the audio, which is a key attribute of comedy rather than action.

Figure 25: A sampled error case in the film and music matching task.

A: the third image
B: the second image
C: the fourth image
D: the first image



[img4]



Error Reason: First, an audio understanding error occurred, as it failed to accurately interpret the audio input, which features slow-paced music with a long, sustained vocal at the end. Instead, the model incorrectly identified the audio as having a moderately paced, syncopated rhythm with a swing feel, entirely mischaracterizing its tempo and articulation. Second, the model incorrectly concluded that the fourth image, with its eighth and sixteenth notes and staccato markings, matches the rhythmic character. However, staccato markings, indicating short, detached articulation, conflict with the described swing feel, which requires smoother, connected phrasing. This reliance on rhythm alone is limiting, as pitch is equally important in musical notation. The model needs to capture the pitch in the audio and align it with the notes in the images to make accurate judgments.

Figure 26: A sampled error case in the music score matching task.

✖

Space: Audio 3D Angle Estimation

Question: Based on [audio1], could you provide the azimuth and elevation angles of the sound located in front of the man dressed in white and black, as illustrated in [img1]? The angles are measured egocentrically, with the azimuth ranging from -180° to 180° (left to right) and the elevation ranging from -90° to 90° (bottom to top).

A: azimuth: 0° , elevation: -1°
 B: azimuth: 38° , elevation: 9°
 C: azimuth: 51° , elevation: -6°
 D: azimuth: 54° , elevation: 10°

[audio1]

audio content: 4-channel spatial audio featuring men's voice and the sound of water flowing from a tap

[img1]

Gemini 1.5 Pro: The man in the maroon shirt is the one dressed in white and black. The microphone pointed at him appears to be at approximately 0° azimuth (directly in front) and about -1° elevation (slightly below the horizontal). So the answer is A.

Error Category: Vision Understanding Error, Audio Understanding Error

Error Reason: First, the statement, "The man in the maroon shirt is the one dressed in white and black," is incorrect. The individual dressed in white and black is not the same as the man in the maroon shirt but is located elsewhere in the scene. This indicates a failure to accurately identify and associate visual details, leading to a misunderstanding of the relationship between these individuals.

Second, It failed to identify the sound of water flowing from the tap, which is visually obscured behind the individual dressed in white and black. This highlights a critical limitation in linking audio cues to visually occluded sources. Additionally, the model was unable to accurately determine the 3D spatial orientation of the tap's water flow.

Correct Answer: C: azimuth: 51° , elevation: -6°

Figure 27: A sampled error case in the audio 3D angle estimation task.

A: 143
B: 137
C: 115
D: 109



Error Category: Audio Understanding Error

Correct Answer: B: 137

39

2106
2107
2108
2109
2110
2111
2112
2113
2114
2115
2116
2117
2118
2119
2120
2121
2122
2123
2124
2125
2126
2127
2128
2129
2130
2131
2132
2133
2134
2135
2136
2137
2138
2139
2140
2141
2142
2143
2144
2145
2146
2147
2148
2149
2150
2151
2152
2153
2154
2155
2156
2157
2158
2159

Time: Audio Time Estimation

Question: From [audio1], what timestamps correspond to the action portrayed in [video1]?

A: start time: 0.00 s, end time: 1.35 s
B: start time: 1.31 s, end time: 2.66 s
C: start time: 2.45 s, end time: 3.15 s
D: start time: 4.30 s, end time: 7.16 s

[audio1]

audio content: the sound of a series of actions

[video1]

Gemini 1.5 Pro: The action portrayed, putting utensils into a drawer, primarily occurs from [00:00:04.30] to [00:00:07.16]. Therefore, the answer is ****D****.

Error Category: **Audio Understanding Error**

Error Reason: Gemini 1.5 Pro failed to accurately identify the timing of the action where a knife is placed into a drawer. Specifically, it misinterpreted the action as occurring in the time range start time: 4.30 s to end time: 7.16 s, instead of the correct range start time: 0.00 s to end time: 1.35 s.

Correct Answer: A: start time: 0.00 s, end time: 1.35 s

Figure 29: A sampled error case in the audio time estimation task.

A: the first audio
B: the third audio
C: the fourth audio
D: the second audio



Error Reason: Gemini 1.5 Pro failed to understand the temporal synchronization between visual and audio elements, leading to errors in aligning events across modalities.

Figure 30: A sampled error case in the audio-visual synchronization task.

✖

Time: Action Sequencing

Question: Based on [audio1], which order most likely represents the actions in [img1], [img2], [img3], and [img4]?

A: the first image --> the second image --> the third image --> the fourth image
 B: the third image --> the fourth image --> the second image --> the first image
 C: the first image --> the fourth image --> the third image --> the second image
 D: the third image --> the fourth image --> the first image --> the second image

[audio1]



audio content: the sound of a series of actions



[img1]



[img2]



[img3]



[img4]

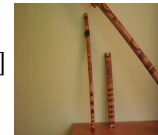
Gemini 1.5 Pro: The correct answer is A. Here's why: **Image 1:** Shows a potato being washed under the tap. This is likely the first step, as raw food preparation usually precedes dishwashing. **Image 2:** Shows a sponge and an empty sink with running water, implying the potato has been washed and now the dishes are being cleaned. **Image 3:** Shows a plate being washed. This follows the cleaning of the potato and the preparation of the sink with soapy water. **Image 4:** Shows the plate being placed in a drying rack with other clean dishes, indicating the dishwashing process is nearing completion. The other options don't make logical sense in terms of the typical sequence of washing a potato and then doing dishes.

Error Category: **Audio Understanding Error, Reasoning Error**
Error Reason: Gemini 1.5 Pro failed to accurately identify the sequence of actions based on the audio cues. The model misinterpreted the order in which sounds corresponding to specific actions occurred, leading to an incorrect understanding of the temporal flow of events. Furthermore, it compounded this error with a reasoning mistake by incorrectly inferring the logical order of the actions.

Correct Answer: D: the third image --> the fourth image --> the first image --> the second image

Figure 31: A sampled error case in the action sequencing task.

A: the second image
B: the third image
C: the fourth image
D: the first image



Error Category: Audio Understanding Error

Correct Answer: A: the second image

43

