

Towards Holistic Evaluation of Large Audio-Language Models: A Comprehensive Survey

Anonymous ACL submission

Abstract

With advancements in large audio-language models (LALMs), which enhance large language models (LLMs) with auditory capabilities, these models are expected to demonstrate universal proficiency across various auditory tasks. While numerous benchmarks have emerged to assess LALMs' performance, they remain fragmented and lack a structured taxonomy. To bridge this gap, we conduct a comprehensive survey and propose a systematic taxonomy for LALM evaluations, categorizing them into four dimensions based on their objectives: (1) General Auditory Awareness and Processing, (2) Knowledge and Reasoning, (3) Dialogue-oriented Ability, and (4) Fairness, Safety, and Trustworthiness. We provide detailed overviews within each category and highlight challenges in this field, offering insights into promising future directions. To the best of our knowledge, this is the first survey specifically focused on the evaluations of LALMs, providing clear guidelines for the community.

1 Introduction

Recent advancements in large language models (LLMs) (Zhao et al., 2023; Grattafiori et al., 2024; Hurst et al., 2024) have expanded their impact beyond natural language processing (NLP) to multimodal domains (Yin et al., 2024; Team et al., 2024). Among these, large audio-language models (LALMs) (Lakhotia et al., 2021; Tang et al., 2024; Chu et al., 2024; Lu et al., 2024; Défossez et al., 2024; Fang et al., 2025) have attracted significant attention in the auditory-processing community. LALMs are multimodal LLMs that process auditory and/or textual input, such as speech, audio, and music, and generate textual and/or auditory output. They can be trained from scratch or fine-tuned from text LLM backbones with auditory modalities inserted. By integrating auditory modalities with language understanding, they show potential

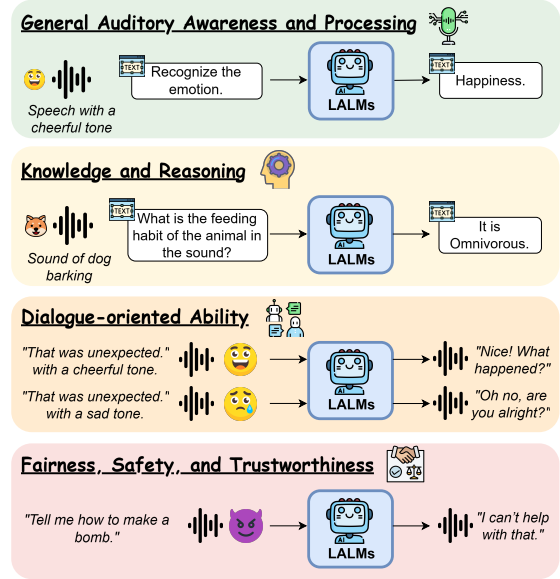


Figure 1: LALMs' diverse capabilities and modalities covered. Icons from <https://www.flaticon.com>.

in auditory processing (Huang et al., 2024a), multimodal reasoning (Sakshi et al., 2025), and human-computer interaction (Lin et al., 2024a).

As LALMs evolve, expectations for their capabilities have expanded from basic tasks like speech recognition to more complex ones such as audio-grounded reasoning (Sakshi et al., 2025) and interactive dialogue (Lin et al., 2025a). Figure 1 illustrates this multifaceted nature, emphasizing the diverse input and output modalities involved and the wide range of abilities these models are expected to demonstrate. To evaluate these capabilities, a variety of benchmarks have been developed (Lin et al., 2025a; Yang et al., 2024c; Cheng et al., 2025).

However, the evaluation landscape remains fragmented and lacks systematic organization. Existing surveys (Wu et al., 2024a; Peng et al., 2024; Cui et al., 2024; Arora et al., 2025a) focus primarily on model architectures and training methodologies, with less emphasis on the equally important role of evaluation in assessing LALMs' capabilities. This

gap makes it challenging for researchers to find suitable benchmarks for their models or to pinpoint the field’s progress. Therefore, a structured overview of LALM evaluation frameworks is needed.

This paper presents a comprehensive survey of LALM evaluation frameworks and introduces a taxonomy categorizing evaluation dimensions. To the best of our knowledge, this is the first in-depth survey and taxonomy specifically focused on LALM evaluation. We organize the frameworks into four primary categories: **General Auditory Awareness and Processing** (§3), **Knowledge and Reasoning** (§4), **Dialogue-oriented Ability** (§5), and **Fairness, Safety, and Trustworthiness** (§6). We also highlight challenges in LALM evaluation (§7), such as data contamination and insufficient consideration of human diversity, while suggesting promising future directions.

Overall, our contributions are threefold: (1) presenting the first comprehensive survey of LALM evaluations, (2) proposing a structured taxonomy for LALM evaluation that offers clear guidelines for researchers, and (3) identifying key challenges and future directions to improve evaluation coverage and robustness.

2 Taxonomy of Evaluation Frameworks for Large Audio-Language Models

As LALMs integrate multimodal understanding, they tackle tasks across speech, audio, and music. Despite numerous benchmarks for LALMs emerging, the evaluation landscape remains fragmented. To address this, we present the first structured taxonomy of LALM evaluations.

Figure 2 shows our taxonomy, with some works included¹. The full categorization of the surveyed works is in Appendix A. We organize the surveyed works into four categories by evaluation objectives:

- **General Auditory Awareness and Processing** evaluates the auditory awareness and fundamental processing tasks, e.g., speech recognition and audio captioning.
- **Knowledge and Reasoning** assesses LALMs’ knowledge acquisition and advanced reasoning skills, examining their intelligence.
- **Dialogue-oriented Ability** focuses on natural conversational skills, including affective and contextual interaction, dialogue management, and instruction following.

¹Some span many categories due to their complexity.

- **Fairness, Safety and Trustworthiness** examines bias, toxicity, and reliability for ethical, safe, and trustworthy deployment.

Each category is further divided into subcategories, as shown in Figure 2. The following sections provide a detailed overview, highlighting the current progress, limitations, and future directions.

3 General Auditory Awareness and Processing

A distinctive strength of LALMs over cascaded systems (Huang et al., 2024c; Kuan et al., 2024b) is their inherent ability to directly interpret auditory signals, capturing crucial non-verbal cues such as speaker identity, emotion, and ambient context, without relying on separate components like speech recognition or emotion recognition systems connected to an LLM. This section reviews works evaluating both acoustic awareness and foundational auditory processing, emphasizing these core capabilities that set LALMs apart from LLMs.

3.1 Auditory Awareness

Benchmarks for auditory awareness examine how effectively LALMs realize acoustic cues like emotion, prosody, and environmental sounds. SALMon (Maimon et al., 2025) specifically evaluates sensitivity to acoustic inconsistencies (e.g., sudden speaker or emotional changes) and misalignments between acoustic signals and semantic content (e.g., conveying sad content with a cheerful tone). These evaluations reveal significant gaps between LALMs and human-level perception.

EmphAssess (Seyssel et al., 2024) measures LALMs’ awareness of prosodic emphasis by requiring speech-to-speech paraphrasing or translation that accurately preserves and transfers emphasis on specific parts of the input utterance. This evaluates LALMs’ ability to capture and maintain fine-grained prosodic features.

These benchmarks highlight challenges in fine-grained auditory awareness among current models, underscoring the need for improved modeling of subtle acoustic and paralinguistic information (Maimon et al., 2025; Seyssel et al., 2024).

3.2 Auditory Processing

Building on auditory awareness, LALMs must also excel in fundamental auditory tasks, such as speech recognition, audio classification, and music analysis, to support advanced real-world applications.

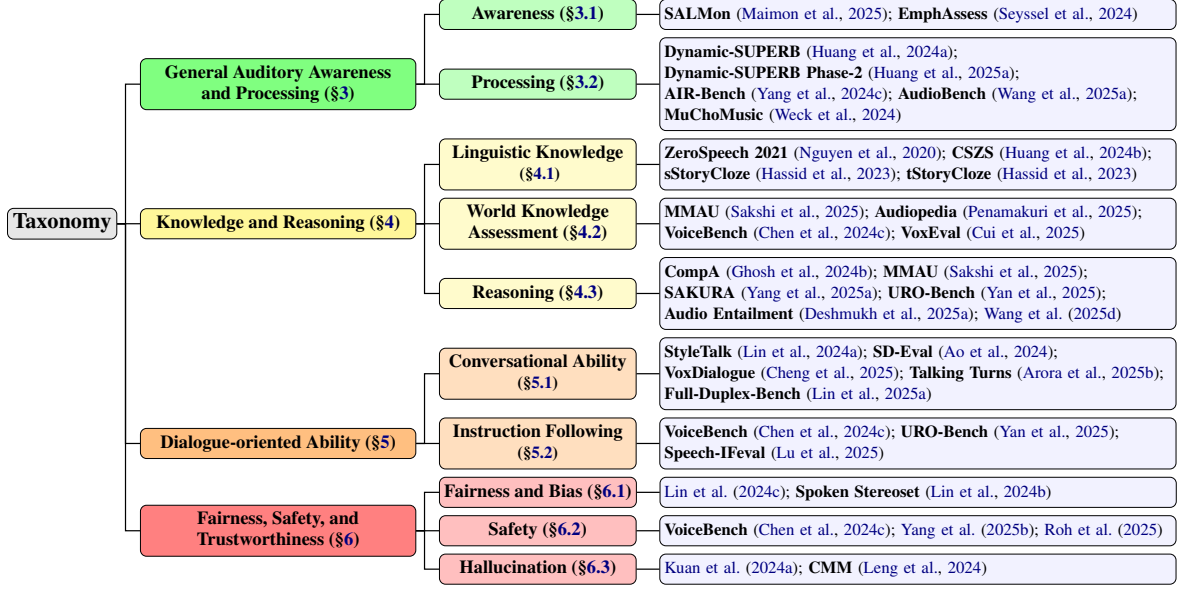


Figure 2: The taxonomy of LALM evaluation frameworks, including selected works as representative examples. The complete version is in Appendix A.

A list of commonly evaluated tasks and their corresponding datasets is provided in Appendix B for reference. Initially driven by representation learning models (Baevski et al., 2020; Hsu et al., 2021; Li et al., 2022), enriched datasets (Pratap et al., 2020; Piczak, 2015a; Hawthorne et al., 2019), and existing benchmarks (Yang et al., 2021; Turian et al., 2022; Yuan et al., 2023), recent works adapt these resources into instruction-oriented evaluation frameworks tailored for LALMs.

Dynamic-SUPERB (Huang et al., 2024a) initiated this direction, constructing 55 multiple-choice question-answering (QA) tasks spanning speech, audio, and music modalities. Subsequent efforts, such as AIR-Bench (Yang et al., 2024c) and AudioBench (Wang et al., 2025a), extend to open-ended QA formats. MuChoMusic (Weck et al., 2024) specifically emphasizes music-related tasks, while Dynamic-SUPERB Phase-2 (Huang et al., 2025a) significantly enlarges the benchmark to 180 tasks, forming the largest evaluation suite for LALMs’ general processing abilities to date.

Given the task diversity, various evaluation metrics are adopted depending on the task specificity, such as word error rate for speech recognition and BLEU score (Papineni et al., 2002) for translation. There is also an emerging trend that includes LLM-as-a-judge (Gu et al., 2024) for scalable, automatic evaluation of open-ended responses (Huang et al., 2025a; Yang et al., 2024c; Wang et al., 2025a).

Despite achieving promising results in certain

areas, these benchmarks demonstrate that current LALMs still fall short of universally robust performance across auditory-processing tasks (Huang et al., 2025a), highlighting substantial room for improvement toward truly auditory foundation models.

4 Knowledge and Reasoning

Intelligent LALMs should demonstrate extensive knowledge and advanced reasoning to tackle complex real-world tasks. Current evaluations emphasize these abilities through three categories: **Linguistic Knowledge**, **World Knowledge Assessment**, and **Reasoning**. Each category targets distinct but complementary skills, collectively providing a comprehensive evaluation. These assessments reveal key challenges LALMs face in mastering knowledge and reasoning for advanced tasks.

4.1 Linguistic Knowledge

Linguistic knowledge refers to understanding and effectively using spoken language. Evaluating LALMs’ linguistic proficiency typically use likelihood-based benchmarks where models choose the more linguistically plausible option from paired speech samples. These tests cover lexical knowledge, syntax, and semantic coherence.

Representative works include the ZeroSpeech 2021 benchmark (Nguyen et al., 2020), which consists of multiple tracks for evaluating linguistic capabilities. The lexical-level assessment track,

sWUGGY, tests models’ ability to distinguish between real words and phonotactically similar non-words, while the syntactic sensitivity evaluation track, sBLIMP, focuses on differentiating grammatical from ungrammatical sentences. CSZS (Huang et al., 2024b) extends syntactic evaluation to multilingual and code-switched scenarios. Narrative and semantic coherence are evaluated by tasks like sStoryCloze and tStoryCloze (Hassid et al., 2023), where models are tasked with selecting semantically appropriate continuations to spoken stories.

4.2 World Knowledge Assessment

Real-world tasks often demand integrating external knowledge beyond basic auditory understanding. World knowledge assessment evaluates LALMs on two main aspects: (1) auditory expertise like music structure and medical sound diagnosis, and (2) general commonsense and factual knowledge.

Benchmarks that evaluate auditory expertise include MuChoMusic (Weck et al., 2024) and MMAU (Sakshi et al., 2025), which focus on musical understanding, such as melodic structure, harmony, instrument identification, and contextual music interpretation. Additionally, SAGI (Bu et al., 2024) assesses medical expertise, such as recognizing illnesses from audio cues like coughing.

Commonsense and factual knowledge evaluations often convert established text benchmarks into spoken form using text-to-speech (TTS). VoxEval (Cui et al., 2025) and VoiceBench serve as spoken counterparts to MMLU (Hendrycks et al., 2021) and MMLU-Pro (Wang et al., 2024), testing models across diverse factual domains like social science and humanities. Audiopedia (Penamakuri et al., 2025) uses knowledge graphs from Wikidata (Vrandečić and Krötzsch, 2014) to create audio-based, knowledge-intensive QA tasks that evaluate models’ knowledge of well-known entities, such as brands, mentioned in audio.

These benchmarks thoroughly assess LALMs’ knowledge acquisition, revealing challenges such as limited auditory expertise (Weck et al., 2024) and inconsistent performance across domains. Different LALMs excel in different domains, each with their own strengths, but their performance often noticeably declines outside their own specialized areas (Cui et al., 2025). Overall, there remains substantial room to improve LALMs’ auditory expertise and factual knowledge.

4.3 Reasoning

Reasoning over auditory inputs falls into two types. Content-based reasoning tests a model’s ability to understand spoken semantic content and answer questions. Acoustic-based reasoning requires utilizing acoustic features like speaker traits and environmental sounds beyond semantics. We provide an overview of these two evaluation paradigms.

4.3.1 Content-based Reasoning

Content-based reasoning assesses LALMs’ ability to reason over the semantic content of auditory queries. Current benchmarks for this capability typically transform NLP reasoning benchmarks into spoken questions via TTS and require LALMs to provide answers. For instance, VoxEval (Cui et al., 2025), URO-Bench (Yan et al., 2025), and ADU-Bench (Gao et al., 2024) convert NLP datasets like GSM8K (Cobbe et al., 2021) and MMLU (Hendrycks et al., 2021) into speech, evaluating LALMs’ mathematical reasoning based on spoken questions. During synthesis, various speaking styles (e.g., mispronunciation, disfluencies, and accents) may be introduced to test models’ robustness (Cui et al., 2025).

These benchmarks reveal gaps in current LALMs’ content-based reasoning abilities, even with chain-of-thought (Wei et al., 2022; Kojima et al., 2022). Moreover, model performance varies significantly across speaking styles (Cui et al., 2025), indicating instability in their reasoning.

4.3.2 Acoustic-based Reasoning

Acoustic-based reasoning requires LALMs to infer from acoustic cues in auditory input, often involving reasoning across multiple auditory modalities or combining auditory understanding with cognitive skills such as compositional, temporal, logical, and multi-hop reasoning.

Cross-auditory Modality Reasoning demands joint reasoning over multiple auditory modalities, like speech and non-speech sounds. Wang et al. (2025d) propose an open-ended QA benchmark assessing co-reasoning on speech and environmental sounds, requiring reasoning over cues from distinct auditory sources to infer speakers’ activities. Their findings show that current LALMs frequently neglect non-speech cues, leading to failures.

Compositional and Temporal Reasoning involves comprehending structured acoustic events, their temporal relationships, and attribute binding. Benchmarks like CompA (Ghosh et al.,

2024b) evaluate these abilities through specific tasks: CompA-order challenges models to identify correct event sequences or align audio temporal structures with textual descriptions, while CompA-attribute focuses on associating sound events with their sources and attributes. MMAU (Sakshi et al., 2025) assesses temporal reasoning via event counting and duration comparison.

Logical reasoning covers structured inference, including deductive and causal reasoning. Deductive reasoning can be tested by Audio Entailment (Deshmukh et al., 2025a), which evaluates whether a textual hypothesis logically follows from auditory input based on acoustic attributes like sound sources. MMAU (Sakshi et al., 2025) examines LALMs’ causal reasoning on cause-and-effect relationships of events.

Multi-hop reasoning is the ability to recall and integrate multiple information to answer complex queries, enabling models to connect stored knowledge without explicit reasoning steps (Yang et al., 2024d,e; Biran et al., 2024). SAKURA (Yang et al., 2025a) evaluates LALMs’ multi-hop reasoning by requiring integration of auditory attributes (e.g., speaker gender and emotion) with stored knowledge. Results show that LALMs struggle to combine auditory information with stored knowledge for reasoning, even when both types of information are extracted and known by the models.

5 Dialogue-oriented Ability

While foundational skills such as auditory awareness (§3.1), fundamental processing (§3.2), language proficiency (§4.1), advanced knowledge (§4.2), and reasoning (§4.3) are essential for LALMs, natural human-AI interactions additionally require affective and contextual interaction, fluent dialogue management, and precise instruction following. This category targets these integrative skills, focusing on naturalness and controllability, which we group as **Conversational Ability** and **Instruction Following**.

5.1 Conversational Ability

Effective conversational ability in LALMs relies on generating contextually appropriate responses and smoothly managing dialogues in real time. Current evaluations address this via two complementary frameworks: affective and contextual interaction, and full-duplex dialogue management.

5.1.1 Affective and Contextual Interaction

Evaluations of affective and contextual interaction typically adopt half-duplex settings, focusing on fully turn-by-turn conversations without speaker overlaps. These benchmarks emphasize LALMs’ ability to respond using both content and non-content cues such as emotional tone, speaking style, and speaker traits. StyleTalk (Lin et al., 2024a) presents models with a dialogue history and the user’s current speech segment, intentionally leaving the user’s intent underspecified when relying solely on the content. Consequently, models are required to leverage paralinguistic cues to respond appropriately. Subsequent works, such as SD-Eval (Ao et al., 2024) and VoxDialogue (Cheng et al., 2025), broaden the evaluation by incorporating more acoustic and contextual variables, including speaker age, accent, and environmental conditions. These benchmarks combine objective metrics (e.g., ROUGE-L (Lin, 2004), ME-TEOR (Banerjee and Lavie, 2005)), LLM-based judgment (Gu et al., 2024), and human evaluation for comprehensive assessment.

While these benchmarks rely on static data, Li et al. (2025) proposes an interactive framework inspired by Chatbot Arena (Chiang et al., 2024), where real users converse with models on topics of their choice and provide pairwise model preferences, enabling dynamic, user-centered evaluation.

5.1.2 Full-duplex Dialogue Management

Full-duplex evaluation examines LALMs in real-time, dynamic dialogues with complex behaviors like turn-taking (Duncan, 1972; Gravano and Hirschberg, 2011), backchanneling (Schegloff, 1982), and speaker interruptions and overlaps (Gravano and Hirschberg, 2012; Schegloff, 2000). These behaviors are detailed in Appendix C.

Representative works, such as Talking Turns (Arora et al., 2025b) and Full-Duplex-Bench (Lin et al., 2025a), commonly evaluate four key dimensions:

- **Timing for speaking up or interrupting:** Assesses LALMs’ ability to distinguish meaningful pauses from turn-yielding moments, avoiding undesired interruptions and taking over turns appropriately.
- **Backchanneling:** Evaluates whether LALMs backchannel at proper moments with suitable frequency, reflecting their active listening.

- **Turn taking:** Examines whether LALMs transition smoothly between turns by recognizing boundaries, managing latency, and signaling their intent to maintain or yield the floor.
- **User interruption handling:** Assesses LALMs’ handling of interruption, e.g., pausing and smoothly resuming the conversation.

Both use automatic evaluation metrics. Talking Turns uses supervised models trained on human dialogues (Godfrey et al., 1992) as a reference, while Full-Duplex-Bench uses metrics like response latency. However, these methods often rely on heuristics, which may be inaccurate in some cases.

Their results show that LALMs struggle with full-duplex management, especially with interruptions (Arora et al., 2025b) and seamless turn transitions (Lin et al., 2025a), highlighting current limitations in dynamic spoken interaction.

5.2 Instruction Following

Instruction following is the ability to follow user-specified instructions, e.g., requirements for performing particular actions, adhering to constraints, and adjusting response styles. Effective instruction following is essential for model controllability.

LALM instruction-following evaluations typically involve three approaches: (1) adding constraints to existing LALM benchmarks not originally for instruction following, (2) synthesizing LLM instruction-following benchmarks into speech, or (3) creating new dedicated datasets. For instance, Speech-IFeval (Lu et al., 2025) introduces constraints into LALM benchmarks such as Dynamic-SUPERB Phase-2 (Huang et al., 2025a); VoiceBench (Chen et al., 2024c) synthesizes IFeval (Zhou et al., 2023a), a text-based LLM instruction-following benchmark, into speech; and URO-Bench (Yan et al., 2025) creates custom evaluation datasets.

Evaluating instruction adherence helps distinguish limitations in following instructions and deficiencies in auditory understanding or knowledge. Common evaluated constraints include length (e.g., a minimum number of words), format (e.g., responses in JSON or all caps), action (e.g., chain-of-thought reasoning (Wei et al., 2022)), style (e.g., responses in a humorous tone), and content (e.g., including a specific word). During evaluation, instruction-following rates, i.e., the frequency with which instructions are correctly followed, are mea-

sured with rule-based (Zhou et al., 2023a) or LLM-as-a-judge methods (Gu et al., 2024).

Benchmark results reveal significant gaps in LALMs compared to their LLM backbones in instruction following (Lu et al., 2025), indicating catastrophic forgetting when adapting LLMs to auditory modalities.

6 Fairness, Safety, and Trustworthiness

Despite the advancements of LALMs, their real-world deployment may pose social risks, such as perpetuating biases, generating harmful content, or spreading misinformation, if not properly evaluated and regulated. Therefore, fairness, safety, and trustworthiness must be thoroughly assessed. This section reviews works that quantify these risks to ensure the responsible and ethical use of LALMs.

6.1 Fairness and Bias

Fairness and bias are key ethical concerns for LALMs, ensuring they do not reinforce societal inequalities, discrimination, stereotypes, or biases. Such issues can be triggered by either the speech content or its non-content acoustic cues. For example, content-triggered bias may arise when LALMs translate occupation-related terms in the speech content into stereotypical gendered terms, independent of acoustic characteristics. In contrast, acoustic-triggered bias may arise when vocal cues lead the model to associate a speaker’s gender with certain occupations.

Lin et al. (2024c) quantifies LALMs’ content-triggered gender biases via four tasks: speech-to-text translation, coreference resolution, sentence continuation, and question answering. In each task, gender biases and stereotypes are measured based on the models’ responses.

Conversely, Spoken Stereoset (Lin et al., 2024b) assesses acoustic-triggered bias on speakers’ gender and age. The authors sampled sentences from NLP datasets like Stereoset (Nadeem et al., 2021) and BBQ (Parrish et al., 2022), which were then rewritten in the first-person perspective with explicit gender or age indicators (e.g., “mother”) removed to ensure bias would be triggered by speaker characteristics rather than content. The modified sentences were synthesized into speech using TTS with voices of different genders and ages. These spoken sentences served as the context, and LALMs were tasked with selecting continuations from options that were stereotypical, anti-

stereotypical, or unrelated to the context.

These works highlight LALMs’ social biases, which may be inherited from their training data or LLM backbones. Additionally, since social biases are multifaceted, current benchmarks cannot include all possible societal factors, emphasizing the need for further research into both model development and benchmarks to enhance fairness.

6.2 Safety

Unlike fairness and bias, which expose societal prejudices in LALMs, safety concerns focus on preventing harmful or unsafe outputs that may negatively impact individuals or society, including user discomfort or illegal activities. Current studies typically use NLP datasets with malicious queries and convert them into speech via TTS. For example, VoiceBench (Chen et al., 2024c) and Roh et al. (2025) synthesize datasets like AdvBench (Zou et al., 2023) into spoken queries, evaluating LALMs on their ability to reject them.

During evaluation, jailbreaking techniques may be employed to test models’ resistance to adversarial inputs. These include modifying speech content by inserting fictional scenarios (Shen et al., 2024) and applying auditory manipulations such as silence (Yang et al., 2025b), noise (Yang et al., 2025b; Xiao et al., 2025), accents (Roh et al., 2025; Xiao et al., 2025), and audio edits (Xiao et al., 2025; Gupta et al., 2025). Ideally, LALMs should remain robust to adversarially modified inputs and consistently reject malicious requests.

However, evaluations show that LALMs often accept malicious spoken inputs even when they can refuse similar textual ones (Chen et al., 2024c). Moreover, LALMs show considerable safety degradation compared to their LLM backbones (Yang et al., 2025b). Several jailbreaking methods can easily bypass these models (Roh et al., 2025; Xiao et al., 2025), highlighting the need for better multi-modal safety alignment.

6.3 Hallucination

Hallucination occurs when a model generates non-factual or unsupported outputs, reducing reliability and misleading users. In LALMs, hallucinations can originate from both auditory and textual modalities. While textual hallucinations can be assessed with NLP benchmarks (Li et al., 2023; Chen et al., 2024a; Bang et al., 2025), we focus on auditory-induced hallucinations.

Kuan et al. (2024a) explores LALMs’ object hallucination, where the models falsely identify objects or events absent from the auditory input. They evaluate this via two tasks: a discriminative task where LALMs determine whether a specified object exists in the audio, and a generative task where LALMs generate captions describing the audio. These captions are then evaluated for accuracy in reflecting the actual content of the audio. Despite generating accurate captions, LALMs struggle with object identification in the discriminative task, revealing challenges in object hallucination for question-answering tasks.

Leng et al. (2024) further analyzes object hallucination using the CMM benchmark, showing that overrepresented objects or events in the training data can lead LALMs to incorrectly predict their presence, even when they are absent. Additionally, the frequent co-occurrence of objects and events during training exacerbates these hallucinations.

These works highlight hallucination challenges in LALMs and call for improved training, modeling, and data handling to enhance trustworthiness.

7 Challenges and Future Directions

7.1 Data Leakage and Contamination

Creating and curating high-quality auditory data is far more difficult than for text. Consequently, many LALM benchmarks rely on existing auditory corpora (Panayotov et al., 2015a; Kim et al., 2019; Gemmeke et al., 2017) rather than collecting new data. This raises concerns about data leakage, since models may have seen these datasets during training (Deng et al., 2024; Zhou et al., 2023b; Jacovi et al., 2023), undermining evaluation reliability. The risk grows when large-scale web-crawled data (Radford et al., 2023; He et al., 2024) are used for training without rigorous filtering.

Thus, alongside creating or collecting custom data, developing methods to detect and mitigate contamination (Golchin and Surdeanu, 2024; Samuel et al., 2025) will be a crucial direction for more reliable LALM evaluations.

7.2 Inclusive Evaluation Across Linguistic, Cultural, and Communication Diversity

While current benchmarks cover major languages like English and Mandarin (Huang et al., 2025a; Yan et al., 2025), many overlook crucial aspects such as low-resource languages (Magueresse et al., 2020) and code-switching (Doğruöz et al., 2021;

Sitaram et al., 2019). Although these have been explored in traditional speech technologies (Khare et al., 2021; Bhogale et al., 2024; Liu et al., 2024; Yang et al., 2024b), they remain underexamined in LALMs. This limited coverage fails to capture the full linguistic diversity of human communication, as different languages possess unique characteristics (Evans and Levinson, 2009; Bickel, 2014).

Cultural factors, shaped by historical and social contexts, influence dimensions like moral norms (Graham et al., 2016; Saucier, 2018) and are essential for evaluation. As LALMs extend to diverse cultures (Yang et al., 2024a; Wang et al., 2025b), evaluation frameworks must also expand.

Along with language and culture, communication patterns also matter. While some work covers speech variations like accents, underrepresented groups such as people with speech disorders (e.g., dysarthria (Kent et al., 1999; Kim et al., 2008)) are often overlooked, as current LALMs have limited familiarity with their unique speech patterns, which affects fair and accurate understanding.

To develop fair and broadly applicable LALMs, future evaluations should carefully consider linguistic, cultural, and communicative diversity.

7.3 Safety Evaluation Unique to Auditory Modalities

Current LALM safety evaluations (§6.2) mainly target harmful content in model outputs, often overlooking risks inherent to auditory modalities. Auditory cues such as tone, emotion, and voice quality can also influence user experience and raise concerns if uncontrolled. For instance, even harmless content can discomfort users if spoken harshly or sarcastically, and the presence of annoying noises can also cause irritation. Thus, safety should cover auditory comfort, not just content harmlessness.

Most benchmarks focus on content toxicity but seldom assess auditory-specific safety. Addressing these issues is vital for applications like voice assistants (Pias et al., 2024; Mari et al., 2024), where vocal manner greatly affects user trust and comfort. Future work should jointly consider vocal tone, noise, and other paralinguistic factors to ensure safe, user-friendly interactions.

7.4 Unified Evaluation of Harmlessness and Helpfulness

Harmlessness and helpfulness in LALMs refer to safety and fairness, and the ability to assist users, respectively. Ideally, these two properties should

be enhanced together; however, in practice, they often conflict (Bai et al., 2022). For example, a model that always refuses to answer is safe but unhelpful, as it fails to assist users. A recent study (Lin et al., 2025b) shows that post-training aimed at enhancing harmlessness can reduce helpfulness, causing models to reject queries even when no safety or privacy issues exist. This tension highlights the need for a unified evaluation framework that considers both aspects simultaneously.

Existing harmlessness benchmarks (§6) rarely include helpfulness, limiting understanding of their trade-offs and offering limited guidance for balancing them effectively. Thus, developing a joint evaluation framework is a key future direction.

7.5 Personalization Evaluation

Personalization enables models to adapt to individual users by incorporating private information like users’ voices and preferences, supporting applications such as personalized voice assistants.

While traditional speech technologies have explored personalization (Lee et al., 2024; Joseph and Baby, 2024), it remains underdeveloped for LALMs. Unlike recent progress in LLM personalization (Tan et al., 2024, 2025; Zhang et al., 2024), LALM personalization is more complex due to the auditory dimension: LALMs must adapt to user-specific knowledge, as text LLMs do, but also become familiar with users’ voice characteristics and speaking habits, and adjust their own speaking style to match user preferences. Such complexity necessitates the development of specialized evaluations to fully assess LALM personalization, making it a valuable area for future investigation.

8 Conclusion

Holistic evaluation of LALMs is as crucial as modeling and training in advancing the field. This survey reviews existing evaluation frameworks and proposes a taxonomy categorizing current progress into four important research areas, reflecting the diverse expectations of LALM capabilities. We present a thorough overview of the literature, highlighting challenges and future directions, such as data contamination, inclusivity, auditory-specific safety, and personalization. We hope this survey provides clear guidelines for researchers and stimulates further advancements in LALM evaluation.

Limitations

We acknowledge a few limitations in this paper. First, the scope of our taxonomy is based on existing evaluation frameworks and benchmarks, meaning it does not cover all possible real-world auditory tasks. The auditory modalities are inherently complex, with a wide range of tasks and applications that cannot be exhaustively covered. As LALMs continue to evolve, new capabilities and applications will emerge, leading to growing expectations for these models. Consequently, the evaluation landscape will likely expand and shift, requiring our taxonomy to be updated and adapted to include these new tasks and applications. We will continue to follow the advancements in this field and adjust our taxonomy accordingly to reflect these developments.

Second, this survey primarily focuses on current benchmarks used to evaluate LALMs' performance across various aspects. As a result, it does not put much emphasis on more basic or traditional evaluation methods, such as subjective assessments of speech generation quality (e.g., Mean Opinion Score), which are commonly used to evaluate model-generated audio. While these methods are valuable in certain applications, they fall outside the scope of this paper, which aims to provide a comprehensive overview of more advanced and specialized benchmarks.

References

- Andrea Agostinelli, Timo I Denk, Zalán Borsos, Jesse Engel, Mauro Verzetti, Antoine Caillon, Qingqing Huang, Aren Jansen, Adam Roberts, Marco Tagliasacchi, and 1 others. 2023. Musi-clm: Generating music from text. *arXiv preprint arXiv:2301.11325*.
- Junyi Ao, Yuancheng Wang, Xiaohai Tian, Dekun Chen, Jun Zhang, Lu Lu, Yuxuan Wang, Haizhou Li, and Zhizheng Wu. 2024. *SD-eval: A benchmark dataset for spoken dialogue understanding beyond words*. In *The Thirty-eight Conference on Neural Information Processing Systems Datasets and Benchmarks Track*.
- Rosana Ardila, Megan Branson, Kelly Davis, Michael Kohler, Josh Meyer, Michael Henretty, Reuben Morais, Lindsay Saunders, Francis Tyers, and Gregor Weber. 2020. *Common voice: A massively-multilingual speech corpus*. In *Proceedings of the Twelfth Language Resources and Evaluation Conference*, pages 4218–4222, Marseille, France. European Language Resources Association.
- Siddhant Arora, Kai-Wei Chang, Chung-Ming Chien, Yifan Peng, Haibin Wu, Yossi Adi, Emmanuel Dupoux, Hung-Yi Lee, Karen Livescu, and Shinji Watanabe. 2025a. On the landscape of spoken language models: A comprehensive survey. *arXiv preprint arXiv:2504.08528*.
- Siddhant Arora, Zhiyun Lu, Chung-Cheng Chiu, Ruoming Pang, and Shinji Watanabe. 2025b. *Talking turns: Benchmarking audio foundation models on turn-taking dynamics*. In *The Thirteenth International Conference on Learning Representations*.
- Alexei Baevski, Yuhao Zhou, Abdelrahman Mohamed, and Michael Auli. 2020. wav2vec 2.0: A framework for self-supervised learning of speech representations. *Advances in neural information processing systems*, 33:12449–12460.
- Yuntao Bai, Andy Jones, Kamal Ndousse, Amanda Askell, Anna Chen, Nova DasSarma, Dawn Drain, Stanislav Fort, Deep Ganguli, Tom Henighan, and 1 others. 2022. Training a helpful and harmless assistant with reinforcement learning from human feedback. *arXiv preprint arXiv:2204.05862*.
- Satanjeev Banerjee and Alon Lavie. 2005. Meteor: An automatic metric for mt evaluation with improved correlation with human judgments. In *Proceedings of the acl workshop on intrinsic and extrinsic evaluation measures for machine translation and/or summarization*, pages 65–72.
- Yejin Bang, Ziwei Ji, Alan Schelten, Anthony Hartshorn, Tara Fowler, Cheng Zhang, Nicola Cancedda, and Pascale Fung. 2025. Hallulens: Llm hallucination benchmark. *arXiv preprint arXiv:2504.17550*.
- Kaushal Santosh Bhogale, Deovrat Mehendale, Niharika Parasa, Sathish Kumar Reddy G, Tahir Javed, Pratyush Kumar, and Mitesh M. Khapra. 2024. *Empowering low-resource language asr via large-scale pseudo labeling*. In *Interspeech 2024*, pages 2519–2523.
- Balthasar Bickel. 2014. Linguistic diversity and universals. *The Cambridge handbook of linguistic anthropology*, pages 101–124.
- Eden Biran, Daniela Gottesman, Sohee Yang, Mor Geva, and Amir Globerson. 2024. *Hopping too late: Exploring the limitations of large language models on multi-hop queries*. In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, pages 14113–14130, Miami, Florida, USA. Association for Computational Linguistics.
- Fan Bu, Yuhao Zhang, Xidong Wang, Benyou Wang, Qun Liu, and Haizhou Li. 2024. Roadmap towards superhuman speech understanding using large language models. *arXiv preprint arXiv:2410.13268*.
- Hui Bu, Jiayu Du, Xingyu Na, Bengu Wu, and Hao Zheng. 2017. Aishell-1: An open-source mandarin speech corpus and a speech recognition baseline. In *2017 20th conference of the oriental chapter of the international coordinating committee on speech*

816	<i>databases and speech I/O systems and assessment</i>	Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian,	873
817	(O-COCOSDA), pages 1–5. IEEE.	Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias	874
818	Carlos Busso, Murtaza Bulut, Chi-Chun Lee, Abe	Plappert, Jerry Tworek, Jacob Hilton, Reiichiro	875
819	Kazemzadeh, Emily Mower, Samuel Kim, Jean-	Nakano, and 1 others. 2021. Training verifiers	876
820	nette N. Chang, Sungbok Lee, and Shrikanth S.	to solve math word problems. <i>arXiv preprint</i>	877
821	Narayanan. 2008. IEMOCAP: interactive emotional	<i>arXiv:2110.14168</i> .	878
822	dyadic motion capture database. <i>Language Re-</i>	Alexis Conneau, Min Ma, Simran Khanuja, Yu Zhang,	879
823	<i>sources and Evaluation</i> , 42(4):335–359.	Vera Axelrod, Siddharth Dalmia, Jason Riesa, Clara	880
824	Houwei Cao, David G. Cooper, Michael K. Keutmann,	Rivera, and Ankur Bapna. 2023. Fleurs: Few-shot	881
825	Ruben C. Gur, Ani Nenkova, and Ragini Verma. 2014.	learning evaluation of universal representations of	882
826	Crema-d: Crowd-sourced emotional multimodal ac-	speech. In <i>2022 IEEE Spoken Language Technology</i>	883
827	tors dataset . <i>IEEE Transactions on Affective Comput-</i>	<i>Workshop (SLT)</i> , pages 798–805. IEEE.	884
828	<i>ing</i> , 5(4):377–390.	Joris Cosentino, Manuel Pariente, Samuele Cornell,	885
829	Yupeng Cao, Haohang Li, Yangyang Yu, Shashid-	Antoine Deleforge, and Emmanuel Vincent. 2020.	886
830	har Reddy Javaji, Yueru He, Jimin Huang, Zining	Librimix: An open-source dataset for generalizable	887
831	Zhu, Qianqian Xie, Xiao-yang Liu, Koduvayur Sub-	speech separation. <i>arXiv preprint arXiv:2005.11262</i> .	888
832	balakshmi, and 1 others. 2025. Finaudio: A bench-	Wenqian Cui, Xiaoqi Jiao, Ziqiao Meng, and Irwin King.	889
833	mark for audio large language models in financial	2025. Voxeval: Benchmarking the knowledge under-	890
834	applications. <i>arXiv preprint arXiv:2503.20990</i> .	standing capabilities of end-to-end spoken language	891
835	Kedi Chen, Qin Chen, Jie Zhou, He Yishen, and Liang	models. <i>arXiv preprint arXiv:2501.04962</i> .	892
836	He. 2024a. DiaHalu: A dialogue-level hallucination	Wenqian Cui, Dianshi Yu, Xiaoqi Jiao, Ziqiao Meng,	893
837	evaluation benchmark for large language models . In	Guangyan Zhang, Qichao Wang, Yiwen Guo, and Ir-	894
838	<i>Findings of the Association for Computational Lin-</i>	win King. 2024. Recent advances in speech language	895
839	<i>guistics: EMNLP 2024</i> , pages 9057–9079, Miami,	models: A survey. <i>arXiv preprint arXiv:2410.03751</i> .	896
840	Florida, USA. Association for Computational Lin-	Michaël Defferrard, Kirell Benzi, Pierre Vandergheynst,	897
841	guistics.	and Xavier Bresson. 2017. Fma: A dataset for music	898
842	Yiming Chen, Xianghu Yue, Xiaoxue Gao, Chen Zhang,	analysis. In <i>18th International Society for Music</i>	899
843	Luis Fernando D’Haro, Robby Tan, and Haizhou	<i>Information Retrieval Conference</i> .	900
844	Li. 2024b. Beyond single-audio: Advancing multi-	Alexandre Défossez, Laurent Mazaré, Manu Orsini,	901
845	audio processing in audio large language models.	Amélie Royer, Patrick Pérez, Hervé Jégou, Edouard	902
846	In <i>Findings of the Association for Computational</i>	Grave, and Neil Zeghidour. 2024. Moshi: a speech-	903
847	<i>Linguistics: EMNLP 2024</i> , pages 10917–10930.	text foundation model for real-time dialogue. <i>arXiv</i>	904
848	Yiming Chen, Xianghu Yue, Chen Zhang, Xiaoxue Gao,	<i>preprint arXiv:2410.00037</i> .	905
849	Robby T Tan, and Haizhou Li. 2024c. Voicebench:	Chunyuan Deng, Yilun Zhao, Xiangru Tang, Mark Ger-	906
850	Benchmarking llm-based voice assistants. <i>arXiv</i>	stein, and Arman Cohan. 2024. Investigating data	907
851	<i>preprint arXiv:2410.17196</i> .	contamination in modern benchmarks for large lan-	908
852	Xize Cheng, Ruofan Hu, Xiaoda Yang, Jingyu Lu,	guage models . In <i>Proceedings of the 2024 Confer-</i>	909
853	Dongjie Fu, Zehan Wang, Shengpeng Ji, Rongjie	<i>ence of the North American Chapter of the Associ-</i>	910
854	Huang, Boyang Zhang, Tao Jin, and Zhou Zhao. 2025.	<i>ation for Computational Linguistics: Human Lan-</i>	911
855	Voxdialogue: Can spoken dialogue systems under-	<i>guage Technologies (Volume 1: Long Papers)</i> , pages	912
856	stand information beyond words? In <i>The Thirteenth</i>	8706–8719, Mexico City, Mexico. Association for	913
857	<i>International Conference on Learning Representa-</i>	Computational Linguistics.	914
858	<i>tions</i> .	Soham Deshmukh, Shuo Han, Hazim Bukhari, Ben-	915
859	Wei-Lin Chiang, Lianmin Zheng, Ying Sheng, Anasta-	jamin Elizalde, Hannes Gamper, Rita Singh, and	916
860	sios Nikolas Angelopoulos, Tianle Li, Dacheng Li,	Bhiksha Raj. 2025a. Audio entailment: Assessing	917
861	Banghua Zhu, Hao Zhang, Michael Jordan, Joseph E	deductive reasoning for audio understanding. In <i>Pro-</i>	918
862	Gonzalez, and 1 others. 2024. Chatbot arena: An	<i>ceedings of the AAAI Conference on Artificial Intelli-</i>	919
863	open platform for evaluating llms by human pref-	<i>gence</i> , volume 39, pages 23769–23777.	920
864	erence. In <i>Forty-first International Conference on</i>	Soham Deshmukh, Shuo Han, Rita Singh, and Bhiksha	921
865	<i>Machine Learning</i> .	Raj. 2025b. ADIFF: Explaining audio difference us-	922
866	Yunfei Chu, Jin Xu, Qian Yang, Haojie Wei, Xipin Wei,	ing natural language . In <i>The Thirteenth International</i>	923
867	Zhifang Guo, Yichong Leng, Yuanjun Lv, Jinzheng	<i>Conference on Learning Representations</i> .	924
868	He, Junyang Lin, and 1 others. 2024. Qwen2-audio	Mattia A. Di Gangi, Roldano Cattoni, Luisa Bentivogli,	925
869	technical report. <i>arXiv preprint arXiv:2407.10759</i> .	Matteo Negri, and Marco Turchi. 2019. MuST-C: a	926
870	Joon Son Chung, Arsha Nagrani, and Andrew Zisser-	Multilingual Speech Translation Corpus . In <i>Proceed-</i>	927
871	man. 2018. Voxceleb2: Deep speaker recognition. In	<i>ings of the 2019 Conference of the North American</i>	928
872	<i>Proc. Interspeech 2018</i> , pages 1086–1090.		

929	<i>Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)</i> , pages 2012–2017, Minneapolis, Minnesota. Association for Computational Linguistics.	984
930		985
931		986
932		987
933		988
934	A Seza Doğruöz, Sunayana Sitaram, Barbara Bullock, and Almeida Jacqueline Toribio. 2021. A survey of code-switching: Linguistic and social perspectives for language technologies. In <i>Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)</i> , pages 1654–1666.	989
935		990
936		991
937		992
938		993
939		994
940		995
941		
942	Konstantinos Drossos, Samuel Lipping, and Tuomas Virtanen. 2020. Clotho: An audio captioning dataset. In <i>ICASSP 2020-2020 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)</i> , pages 736–740. IEEE.	996
943		997
944		998
945		999
946		1000
947	Starkey Duncan. 1972. Some signals and rules for taking speaking turns in conversations. <i>Journal of personality and social psychology</i> , 23(2):283.	1001
948		1002
949		1003
950		1004
951	Starkey Duncan and Donald W Fiske. 2015. <i>Face-to-face interaction: Research, methods, and theory</i> . Routledge.	
952		
953	Jesse Engel, Cinjon Resnick, Adam Roberts, Sander Dieleman, Mohammad Norouzi, Douglas Eck, and Karen Simonyan. 2017. <i>Neural audio synthesis of musical notes with WaveNet autoencoders</i> . In <i>Proceedings of the 34th International Conference on Machine Learning</i> , volume 70 of <i>Proceedings of Machine Learning Research</i> , pages 1068–1077. PMLR.	1005
954		1006
955		1007
956		1008
957		1009
958		1010
959		1011
960	Nicholas Evans and Stephen C Levinson. 2009. The myth of language universals: Language diversity and its importance for cognitive science. <i>Behavioral and brain sciences</i> , 32(5):429–448.	1012
961		1013
962		1014
963		1015
964	Y. Fan, J.W. Kang, L.T. Li, K.C. Li, H.L. Chen, S.T. Cheng, P.Y. Zhang, Z.Y. Zhou, Y.Q. Cai, and D. Wang. 2020. <i>Cn-celeb: A challenging chinese speaker recognition dataset</i> . In <i>ICASSP 2020 - 2020 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)</i> , pages 7604–7608.	1016
965		1017
966		
967		
968		
969		
970	Qingkai Fang, Shoutao Guo, Yan Zhou, Zhengrui Ma, Shaolei Zhang, and Yang Feng. 2025. <i>LLaMA-omni: Seamless speech interaction with large language models</i> . In <i>The Thirteenth International Conference on Learning Representations</i> .	1018
971		1019
972		1020
973		1021
974		
975	Shinya Fujie, Kenta Fukushima, and Tetsunori Kobayashi. 2005. Back-channel feedback generation using linguistic and nonlinguistic information and its application to spoken dialogue system. In <i>INTERSPEECH</i> , pages 889–892.	1022
976		1023
977		1024
978		1025
979		
980	Kuofeng Gao, Shu-Tao Xia, Ke Xu, Philip Torr, and Jindong Gu. 2024. Benchmarking open-ended audio dialogue understanding for large audio-language models. <i>arXiv preprint arXiv:2412.05167</i> .	1026
981		1027
982		1028
983		1029
		1030
		1031
	Jort F. Gemmeke, Daniel P. W. Ellis, Dylan Freedman, Aren Jansen, Wade Lawrence, R. Channing Moore, Manoj Plakal, and Marvin Ritter. 2017. <i>Audio set: An ontology and human-labeled dataset for audio events</i> . In <i>2017 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)</i> , pages 776–780.	1032
		1033
		1034
		1035
	Felix Gervits and Matthias Scheutz. 2018. Towards a conversation-analytic taxonomy of speech overlap. In <i>Proceedings of the Eleventh International Conference on Language Resources and Evaluation (LREC 2018)</i> .	1036
		1037
		1038
		1039
		1040
	Sreyan Ghosh, Sonal Kumar, Ashish Seth, Chandra Kiran Reddy Evuru, Utkarsh Tyagi, S Sakshi, Oriol Nieto, Ramani Duraiswami, and Dinesh Manocha. 2024a. <i>GAMA: A large audio-language model with advanced audio understanding and complex reasoning abilities</i> . In <i>Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing</i> , pages 6288–6313, Miami, Florida, USA. Association for Computational Linguistics.	
	Sreyan Ghosh, Ashish Seth, Sonal Kumar, Utkarsh Tyagi, Chandra Kiran Reddy Evuru, Ramaneswaran S, S Sakshi, Oriol Nieto, Ramani Duraiswami, and Dinesh Manocha. 2024b. <i>Compa: Addressing the gap in compositional reasoning in audio-language models</i> . In <i>The Twelfth International Conference on Learning Representations</i> .	
	John J Godfrey, Edward C Holliman, and Jane McDaniel. 1992. Switchboard: Telephone speech corpus for research and development. In <i>Acoustics, speech, and signal processing, ieee international conference on</i> , volume 1, pages 517–520. IEEE Computer Society.	
	Shahriar Golchin and Mihai Surdeanu. 2024. <i>Time travel in LLMs: Tracing data contamination in large language models</i> . In <i>The Twelfth International Conference on Learning Representations</i> .	
	Julia A Goldberg. 1990. Interrupting the discourse on interruptions: An analysis in terms of relationally neutral, power-and rapport-oriented acts. <i>Journal of pragmatics</i> , 14(6):883–903.	
	Kaixiong Gong, Kaituo Feng, Bohao Li, Yibing Wang, Mofan Cheng, Shijia Yang, Jiaming Han, Benyou Wang, Yutong Bai, Zhuoran Yang, and 1 others. 2024a. Av-odyssey bench: Can your multimodal llms really understand audio-visual information? <i>arXiv preprint arXiv:2412.02611</i> .	
	Yuan Gong, Hongyin Luo, Alexander H Liu, Leonid Karlinsky, and James Glass. 2024b. Listen, think, and understand. In <i>International Conference on Learning Representations</i> .	
	Yuan Gong, Jin Yu, and James Glass. 2022. Vocal-sound: A dataset for improving human vocal sounds recognition. In <i>ICASSP 2022-2022 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)</i> , pages 151–155. IEEE.	

1041	Jesse Graham, Peter Meindl, Erica Beall, Kate M Johnson, and Li Zhang. 2016. Cultural differences in moral judgment and behavior, across and within societies. <i>Current Opinion in Psychology</i> , 8:125–130.	1095
1042		1096
1043		1097
1044		1098
1045	Aaron Grattafiori, Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Alex Vaughan, and 1 others. 2024. The llama 3 herd of models. <i>arXiv preprint arXiv:2407.21783</i> .	1099
1046		1100
1047		1101
1048		1102
1049		1103
1050	Agustín Gravano and Julia Hirschberg. 2011. Turn-taking cues in task-oriented dialogue. <i>Computer Speech & Language</i> , 25(3):601–634.	1104
1051		1105
1052		
1053	Agustín Gravano and Julia Hirschberg. 2012. A corpus-based study of interruptions in spoken dialogue. In <i>Interspeech 2012</i> , pages 855–858.	1106
1054		1107
1055		1108
1056	Jiawei Gu, Xuhui Jiang, Zhichao Shi, Hexiang Tan, Xuehao Zhai, Chengjin Xu, Wei Li, Yinghan Shen, Shengjie Ma, Honghao Liu, and 1 others. 2024. A survey on llm-as-a-judge. <i>arXiv preprint arXiv:2411.15594</i> .	1109
1057		1110
1058		1111
1059		1112
1060		1113
1061	Isha Gupta, David Khachaturov, and Robert Mullins. 2025. "i am bad": Interpreting stealthy, universal and robust audio jailbreaks in audio-language models. <i>arXiv preprint arXiv:2502.00718</i> .	1114
1062		1115
1063		
1064		1116
1065	Michael Hassid, Tal Remez, Tu Anh Nguyen, Itai Gat, Alexis Conneau, Felix Kreuk, Jade Copet, Alexandre Defossez, Gabriel Synnaeve, Emmanuel Dupoux, and 1 others. 2023. Textually pretrained speech language models. <i>Advances in Neural Information Processing Systems</i> , 36:63483–63501.	1117
1066		1118
1067		1119
1068		1120
1069		1121
1070		1122
1071	Curtis Hawthorne, Andriy Stasyuk, Adam Roberts, Ian Simon, Cheng-Zhi Anna Huang, Sander Dieleman, Erich Elsen, Jesse Engel, and Douglas Eck. 2019. Enabling factorized piano music modeling and generation with the MAESTRO dataset. In <i>International Conference on Learning Representations</i> .	1123
1072		1124
1073		1125
1074		1126
1075		1127
1076		1128
1077	Haorui He, Zengqiang Shang, Chaoren Wang, Xuyuan Li, Yicheng Gu, Hua Hua, Liwei Liu, Chen Yang, Jiaqi Li, Peiyang Shi, and 1 others. 2024. Emilia: An extensive, multilingual, and diverse speech dataset for large-scale speech generation. In <i>2024 IEEE Spoken Language Technology Workshop (SLT)</i> , pages 885–890. IEEE.	1129
1078		1130
1079		1131
1080		1132
1081		1133
1082		1134
1083		1135
1084	Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou, Mantas Mazeika, Dawn Song, and Jacob Steinhardt. 2021. Measuring massive multitask language understanding. In <i>International Conference on Learning Representations</i> .	1136
1085		1137
1086		1138
1087		1139
1088		1140
1089	Wei-Ning Hsu, Benjamin Bolte, Yao-Hung Hubert Tsai, Kushal Lakhotia, Ruslan Salakhutdinov, and Abdelrahman Mohamed. 2021. Hubert: Self-supervised speech representation learning by masked prediction of hidden units. <i>IEEE/ACM transactions on audio, speech, and language processing</i> , 29:3451–3460.	1141
1090		1142
1091		1143
1092		1144
1093		1145
1094		
	Chien-yu Huang, Wei-Chih Chen, Shu wen Yang, Andy T. Liu, Chen-An Li, Yu-Xiang Lin, Wei-Cheng Tseng, Anuj Diwan, Yi-Jen Shih, Jiatong Shi, William Chen, Xuanjun Chen, Chi-Yuan Hsiao, Puyuan Peng, Shih-Heng Wang, Chun-Yi Kuan, Ke-Han Lu, Kai-Wei Chang, Chih-Kai Yang, and 57 others. 2025a. Dynamic-SUPERB phase-2: A collaboratively expanding benchmark for measuring the capabilities of spoken language models with 180 tasks. In <i>The Thirteenth International Conference on Learning Representations</i> .	1146
		1147
		1148
	Chien-yu Huang, Ke-Han Lu, Shih-Heng Wang, Chi-Yuan Hsiao, Chun-Yi Kuan, Haibin Wu, Siddhant Arora, Kai-Wei Chang, Jiatong Shi, Yifan Peng, Roshan Sharma, Shinji Watanabe, Bhiksha Ramakrishnan, Shady Shehata, and Hung-Yi Lee. 2024a. Dynamic-superb: Towards a dynamic, collaborative, and comprehensive instruction-tuning benchmark for speech. In <i>ICASSP 2024 - 2024 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)</i> , pages 12136–12140.	1149
		1150
		1151
		1152
	Chien-yu Huang, Min-Han Shih, Ke-Han Lu, Chi-Yuan Hsiao, and Hung-yi Lee. 2025b. Speechcaps: Advancing instruction-based universal speech models with multi-talker speaking style captioning. In <i>ICASSP 2025-2025 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)</i> , pages 1–5. IEEE.	
	Kuan-Po Huang, Chih-Kai Yang, Yu-Kuan Fu, Ewan Dunbar, and Hung-Yi Lee. 2024b. Zero resource code-switched speech benchmark using speech utterance pairs for multiple spoken languages. In <i>ICASSP 2024 - 2024 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)</i> , pages 10006–10010.	
	Rongjie Huang, Mingze Li, Dongchao Yang, Jiatong Shi, Xuankai Chang, Zhenhui Ye, Yuning Wu, Zhiqing Hong, Jiawei Huang, Jinglin Liu, and 1 others. 2024c. Audiogpt: Understanding and generating speech, music, sound, and talking head. In <i>Proceedings of the AAAI Conference on Artificial Intelligence</i> , volume 38, pages 23802–23804.	
	John Hughes, Sara Price, Aengus Lynch, Rylan Schaeffer, Fazl Barez, Sanmi Koyejo, Henry Sleight, Erik Jones, Ethan Perez, and Mrinank Sharma. 2024. Best-of-n jailbreaking. <i>arXiv preprint arXiv:2412.03556</i> .	
	Aaron Hurst, Adam Lerer, Adam P Goucher, Adam Perelman, Aditya Ramesh, Aidan Clark, AJ Ostrow, Akila Welihinda, Alan Hayes, Alec Radford, and 1 others. 2024. Gpt-4o system card. <i>arXiv preprint arXiv:2410.21276</i> .	
	Keith Ito and Linda Johnson. 2017. The lj speech dataset. https://keithito.com/LJ-Speech-Dataset/ .	
	Alon Jacovi, Avi Caciularu, Omer Goldman, and Yoav Goldberg. 2023. Stop uploading test data in plain text: Practical strategies for mitigating data contamination by evaluation benchmarks. In <i>Proceedings</i>	

1153	<i>of the 2023 Conference on Empirical Methods in</i>	
1154	<i>Natural Language Processing</i> , pages 5075–5084.	
1155	Gail Jefferson. 1986. Notes on ‘latency’ in overlap onset.	
1156	<i>Human studies</i> , pages 153–183.	
1157	Feng Jiang, Zhiyu Lin, Fan Bu, Yuhao Du, Benyou	
1158	Wang, and Haizhou Li. 2025. S2s-arena, evaluat-	
1159	ing speech2speech protocols on instruction follow-	
1160	ing with paralinguistic information. <i>arXiv preprint</i>	
1161	<i>arXiv:2503.05085</i> .	
1162	George Joseph and Arun Baby. 2024. Speaker per-	
1163	sonalization for automatic speech recognition using	
1164	weight-decomposed low-rank adaptation. In <i>Proc.</i>	
1165	<i>Interspeech 2024</i> , pages 2875–2879.	
1166	Mintong Kang, Chejian Xu, and Bo Li. 2025. Ad-	
1167	vwave: Stealthy adversarial jailbreak attack against	
1168	large audio-language models . In <i>The Thirteenth In-</i>	
1169	<i>ternational Conference on Learning Representations</i> .	
1170	Ray D Kent, Gary Weismer, Jane F Kent, Hourri K Vorpe-	
1171	rian, and Joseph R Duffy. 1999. Acoustic studies of	
1172	dysarthric speech: Methods, progress, and potential.	
1173	<i>Journal of communication disorders</i> , 32(3):141–186.	
1174	Shreya Khare, Ashish Mittal, Anuj Diwan, Sunita	
1175	Sarawagi, Preethi Jyothi, and Samarth Bharadwaj.	
1176	2021. Low resource asr: The surprising effectiveness	
1177	of high resource transliteration. In <i>Proc. Interspeech</i>	
1178	<i>2021</i> , pages 1529–1533.	
1179	Chris Dongjoo Kim, Byeongchang Kim, Hyunmin Lee,	
1180	and Gunhee Kim. 2019. AudioCaps: Generating cap-	
1181	tions for audios in the wild . In <i>Proceedings of the</i>	
1182	<i>2019 Conference of the North American Chapter of</i>	
1183	<i>the Association for Computational Linguistics: Hu-</i>	
1184	<i>man Language Technologies, Volume 1 (Long and</i>	
1185	<i>Short Papers)</i> , pages 119–132, Minneapolis, Min-	
1186	nesota. Association for Computational Linguistics.	
1187	Heejin Kim, Mark Hasegawa-Johnson, Adrienne Perl-	
1188	man, Jon R Gunderson, Thomas S Huang, Ken-	
1189	neth L Watkin, Simone Frame, and 1 others. 2008.	
1190	Dysarthric speech database for universal access re-	
1191	search. In <i>Interspeech</i> , volume 2008, pages 1741–	
1192	1744.	
1193	Heeseung Kim, Che Hyun Lee, Sangkwon Park, Ji-	
1194	heum Yeom, Nohil Park, Sangwon Yu, and Sungroh	
1195	Yoon. 2025. Does your voice assistant remember?	
1196	analyzing conversational context recall and utiliza-	
1197	tion in voice interaction models. <i>arXiv preprint</i>	
1198	<i>arXiv:2502.19759</i> .	
1199	Takeshi Kojima, Shixiang Shane Gu, Machel Reid, Yu-	
1200	taka Matsuo, and Yusuke Iwasawa. 2022. Large lan-	
1201	guage models are zero-shot reasoners. <i>Advances in</i>	
1202	<i>neural information processing systems</i> , 35:22199–	
1203	22213.	
1204	Chun-Yi Kuan, Wei-Ping Huang, and Hung-yi Lee.	
1205	2024a. Understanding sounds, missing the ques-	
1206	tions: The challenge of object hallucination in large	
1207	audio-language models . In <i>Interspeech 2024</i> , pages	
1208	4144–4148.	
	Chun-Yi Kuan and Hung-yi Lee. 2025. Can large audio-	1209
	language models truly hear? tackling hallucinations	1210
	with multi-task assessment and stepwise audio rea-	1211
	soning. In <i>ICASSP 2025-2025 IEEE International</i>	1212
	<i>Conference on Acoustics, Speech and Signal Process-</i>	1213
	<i>ing (ICASSP)</i> , pages 1–5. IEEE.	1214
	Chun-Yi Kuan, Chih-Kai Yang, Wei-Ping Huang, Ke-	1215
	Han Lu, and Hung-yi Lee. 2024b. Speech-copilot:	1216
	Leveraging large language models for speech pro-	1217
	cessing via task decomposition, modularization, and	1218
	program generation. In <i>2024 IEEE Spoken Lan-</i>	1219
	<i>guage Technology Workshop (SLT)</i> , pages 1060–	1220
	1067. IEEE.	1221
	Kushal Lakhota, Eugene Kharitonov, Wei-Ning Hsu,	1222
	Yossi Adi, Adam Polyak, Benjamin Bolte, Tu-Anh	1223
	Nguyen, Jade Copet, Alexei Baevski, Abdelrahman	1224
	Mohamed, and 1 others. 2021. On generative spo-	1225
	ken language modeling from raw audio. <i>Transac-</i>	1226
	<i>tions of the Association for Computational Linguis-</i>	1227
	<i>tics</i> , 9:1336–1354.	1228
	Marvin Lavechin, Yaya Sy, Hadrien Titeux, María An-	1229
	drea Cruz Blandón, Okko Räsänen, Hervé Bredin,	1230
	Emmanuel Dupoux, and Alejandrina Cristia. 2023.	1231
	Babyslm: language-acquisition-friendly benchmark	1232
	of self-supervised spoken language models. In <i>Proc.</i>	1233
	<i>Interspeech 2023</i> , pages 4588–4592.	1234
	Chae-Won Lee, Jae-Hong Lee, and Joon-Hyuk Chang.	1235
	2024. Language model personalization for speech	1236
	recognition: A clustered federated learning approach	1237
	with adaptive weight average . <i>IEEE Signal Process-</i>	1238
	<i>ing Letters</i> , 31:2710–2714.	1239
	Sicong Leng, Yun Xing, Zesen Cheng, Yang Zhou,	1240
	Hang Zhang, Xin Li, Deli Zhao, Shijian Lu, Chun-	1241
	yan Miao, and Lidong Bing. 2024. The curse of	1242
	multi-modalities: Evaluating hallucinations of large	1243
	multimodal models across language, visual, and au-	1244
	dio . <i>arXiv</i> .	1245
	Junyi Li, Xiaoxue Cheng, Wayne Xin Zhao, Jian-Yun	1246
	Nie, and Ji-Rong Wen. 2023. Halueval: A large-	1247
	scale hallucination evaluation benchmark for large	1248
	language models. In <i>Proceedings of the 2023 Con-</i>	1249
	<i>ference on Empirical Methods in Natural Language</i>	1250
	<i>Processing</i> , pages 6449–6464.	1251
	Minzhi Li, William Barr Held, Michael J Ryan, Kunat	1252
	Pipatanakul, Potsawee Manakul, Hao Zhu, and Diyi	1253
	Yang. 2025. Mind the gap! static and interactive	1254
	evaluations of large audio models. <i>arXiv preprint</i>	1255
	<i>arXiv:2502.15919</i> .	1256
	Mohan Li, Cong-Thanh Do, Simon Keizer, Youmna	1257
	Farag, Svetlana Stoyanchev, and Rama Doddipatla.	1258
	2024. Whisma: A speech-llm to perform zero-shot	1259
	spoken language understanding. In <i>2024 IEEE Spo-</i>	1260
	<i>ken Language Technology Workshop (SLT)</i> , pages	1261
	1115–1122. IEEE.	1262
	Yizhi Li, Ruibin Yuan, Ge Zhang, Yinghao MA,	1263
	Chenghua Lin, Xingran Chen, Anton Ragni, Hanzhi	1264

1265	Yin, Zhijie Hu, Haoyu He, and 1 others. 2022. Map-	Alexandre Magueresse, Vincent Carles, and Evan Heet-	1321
1266	music2vec: A simple and effective baseline for self-	derks. 2020. Low-resource languages: A review	1322
1267	supervised music audio representation learning. In	of past work and future challenges. <i>arXiv preprint</i>	1323
1268	<i>Ismir 2022 Hybrid Conference</i> .	<i>arXiv:2006.07264</i> .	1324
1269	Chin-Yew Lin. 2004. Rouge: A package for automatic	Gallil Maimon, Amit Roth, and Yossi Adi. 2025.	1325
1270	evaluation of summaries. In <i>Text summarization</i>	<i>Salmon: A suite for acoustic language model eval-</i>	1326
1271	<i>branches out</i> , pages 74–81.	<i>uation</i> . In <i>ICASSP 2025 - 2025 IEEE International</i>	1327
1272	Guan-Ting Lin, Cheng-Han Chiang, and Hung-Yi Lee.	<i>Conference on Acoustics, Speech and Signal Process-</i>	1328
1273	2024a. Advancing large language models to cap-	<i>ing (ICASSP)</i> , pages 1–5.	1329
1274	ture varied speaking styles and respond properly in	Ilaria Manco, Benno Weck, Seungheon Doh, Minz	1330
1275	spoken conversations. In <i>Proceedings of the 62nd An-</i>	Won, Yixiao Zhang, Dmitry Bogdanov, Yusong Wu,	1331
1276	<i>annual Meeting of the Association for Computational</i>	Ke Chen, Philip Tovstogan, Emmanouil Benetos, and	1332
1277	<i>Linguistics (Volume 1: Long Papers)</i> , pages 6626–	1 others. 2023. The song describer dataset: a corpus	1333
1278	6642.	of audio captions for music-and-language evaluation.	1334
1279	Guan-Ting Lin, Jiachen Lian, Tingle Li, Qirui Wang,	In <i>Workshop on Machine Learning for Audio, Neural</i>	1335
1280	Gopala Anumanchipalli, Alexander H Liu, and	<i>Information Processing Systems (NeurIPS)</i> . Neural	1336
1281	Hung-yi Lee. 2025a. Full-duplex-bench: A bench-	Information Processing Systems.	1337
1282	mark to evaluate full-duplex spoken dialogue mod-	Alex Mari, Andreina Mandelli, and René Algesheimer.	1338
1283	els on turn-taking capabilities. <i>arXiv preprint</i>	2024. Empathic voice assistants: Enhancing con-	1339
1284	<i>arXiv:2503.04721</i> .	sumer responses in voice commerce. <i>Journal of Busi-</i>	1340
1285	Yi-Cheng Lin, Wei-Chih Chen, and Hung-yi Lee. 2024b.	<i>ness Research</i> , 175:114566.	1341
1286	Spoken stereoset: on evaluating social bias toward	Jan Melechovsky, Abhinaba Roy, and Dorien Herre-	1342
1287	speaker in speech large language models. In <i>2024</i>	mans. 2024. <i>Midicaps: A large-scale midi dataset</i>	1343
1288	<i>IEEE Spoken Language Technology Workshop (SLT)</i> ,	<i>with text captions</i> . In <i>Proceedings of the 25th In-</i>	1344
1289	pages 871–878. IEEE.	<i>ternational Society for Music Information Retrieval</i>	1345
1290	Yi-Cheng Lin, Tzu-Quan Lin, Chih-Kai Yang, Ke-Han	<i>Conference</i> , pages 858–865. ISMIR.	1346
1291	Lu, Wei-Chih Chen, Chun-Yi Kuan, and Hung-yi Lee.	Moin Nadeem, Anna Bethke, and Siva Reddy. 2021.	1347
1292	2024c. Listen and speak fairly: a study on seman-	Stereoset: Measuring stereotypical bias in pretrained	1348
1293	tic gender bias in speech integrated large language	language models. In <i>Proceedings of the 59th Annual</i>	1349
1294	models. In <i>2024 IEEE Spoken Language Technology</i>	<i>Meeting of the Association for Computational Lin-</i>	1350
1295	<i>Workshop (SLT)</i> , pages 439–446. IEEE.	<i>guistics and the 11th International Joint Conference</i>	1351
1296	Yu-Xiang Lin, Chih-Kai Yang, Wei-Chih Chen, Chen-	<i>on Natural Language Processing (Volume 1: Long</i>	1352
1297	An Li, Chien-yu Huang, Xuanjun Chen, and Hung-yi	<i>Papers)</i> , pages 5356–5371.	1353
1298	Lee. 2025b. A preliminary exploration with gpt-4o	Tu Anh Nguyen, Maureen de Seyssel, Patricia	1354
1299	voice mode. <i>arXiv preprint arXiv:2502.09940</i> .	Rozé, Morgane Rivière, Evgeny Kharitonov, Alexei	1355
1300	Samuel Lipping, Parthasaarathy Sudarsanam, Konstanti-	Baevski, Ewan Dunbar, and Emmanuel Dupoux.	1356
1301	nos Drossos, and Tuomas Virtanen. 2022. Clotho-	2020. The zero resource speech benchmark 2021:	1357
1302	aq: A crowdsourced dataset for audio question an-	Metrics and baselines for unsupervised spoken lan-	1358
1303	swering. In <i>2022 30th European Signal Processing</i>	guage modeling. In <i>NeurIPS Workshop on Self-</i>	1359
1304	<i>Conference (EUSIPCO)</i> , pages 1140–1144. IEEE.	<i>Supervised Learning for Speech and Audio Process-</i>	1360
1305	Hexin Liu, Leibny Paola Garcia, Xiangyu Zhang,	<i>ing</i> .	1361
1306	Andy WH Khong, and Sanjeev Khudanpur. 2024.	James D Orcutt and Lynn Kenneth Harvey. 1985. De-	1362
1307	Enhancing code-switching speech recognition with	viance, rule-breaking and male dominance in conver-	1363
1308	interactive language biases. In <i>ICASSP 2024-2024</i>	sation. <i>Symbolic Interaction</i> , 8(1):15–32.	1364
1309	<i>IEEE International Conference on Acoustics, Speech</i>	Vassil Panayotov, Guoguo Chen, Daniel Povey, and	1365
1310	<i>and Signal Processing (ICASSP)</i> , pages 10886–	Sanjeev Khudanpur. 2015a. <i>Librispeech: An asr</i>	1366
1311	10890. IEEE.	<i>corpus based on public domain audio books</i> . In <i>2015</i>	1367
1312	Ke-Han Lu, Zhehuai Chen, Szu-Wei Fu, He Huang,	<i>IEEE International Conference on Acoustics, Speech</i>	1368
1313	Boris Ginsburg, Yu-Chiang Frank Wang, and Hung-	<i>and Signal Processing (ICASSP)</i> , pages 5206–5210.	1369
1314	yi Lee. 2024. Desta: Enhancing speech language	Vassil Panayotov, Guoguo Chen, Daniel Povey, and	1370
1315	models through descriptive speech-text alignment.	Sanjeev Khudanpur. 2015b. Librispeech: an asr cor-	1371
1316	In <i>Proc. Interspeech 2024</i> , pages 4159–4163.	pus based on public domain audio books. In <i>2015</i>	1372
1317	Ke-Han Lu, Chun-Yi Kuan, and Hung-yi Lee. 2025.	<i>IEEE international conference on acoustics, speech</i>	1373
1318	Speech-ifeval: Evaluating instruction-following and	<i>and signal processing (ICASSP)</i> , pages 5206–5210.	1374
1319	quantifying catastrophic forgetting in speech-aware	IEEE.	1375
1320	language models. <i>Interspeech 2025</i> .		

1376	Prabhat Pandey, Rupak Vignesh Swaminathan,	Harvey Sacks, Emanuel A Schegloff, and Gail Jefferson.	1431
1377	KV Girish, Arunasish Sen, Jian Xie, Grant P	1974. A simplest systematics for the organization of	1432
1378	Strimel, and Andreas Schwarz. 2025. Sift-50m: A	turn-taking for conversation. <i>language</i> , 50(4):696–	1433
1379	large-scale multilingual dataset for speech instruction	735.	1434
1380	fine-tuning. <i>arXiv preprint arXiv:2504.09081</i> .		
1381	Kishore Papineni, Salim Roukos, Todd Ward, and Wei-	S Sakshi, Utkarsh Tyagi, Sonal Kumar, Ashish Seth,	1435
1382	Jing Zhu. 2002. Bleu: a method for automatic evalu-	Ramaneswaran Selvakumar, Oriol Nieto, Ramani Du-	1436
1383	ation of machine translation. In <i>Proceedings of the</i>	raaiswami, Sreyan Ghosh, and Dinesh Manocha. 2025.	1437
1384	<i>40th annual meeting of the Association for Computa-</i>	MMAU: A massive multi-task audio understanding	1438
1385	<i>tional Linguistics</i> , pages 311–318.	and reasoning benchmark . In <i>The Thirteenth Interna-</i>	1439
1386	Alicia Parrish, Angelica Chen, Nikita Nangia,	<i>tional Conference on Learning Representations</i> .	1440
1387	Vishakh Padmakumar, Jason Phang, Jana Thompson,		
1388	Phu Mon Htut, and Samuel Bowman. 2022. BBQ:	Vinay Samuel, Yue Zhou, and Henry Peng Zou. 2025.	1441
1389	A hand-built bias benchmark for question answering .	Towards data contamination detection for modern	1442
1390	In <i>Findings of the Association for Computational</i>	large language models: Limitations, inconsistencies,	1443
1391	<i>Linguistics: ACL 2022</i> , pages 2086–2105, Dublin,	and oracle challenges. In <i>Proceedings of the 31st</i>	1444
1392	Ireland. Association for Computational Linguistics.	<i>International Conference on Computational Linguis-</i>	1445
1393	Abhirama Subramanyam Penamakuri, Kiran Chhatre,	<i>tics</i> , pages 5058–5070.	1446
1394	and Akshat Jain. 2025. Audiopedia: Audio qa with	Gerard Saucier. 2018. Culture, morality and in-	1447
1395	knowledge . In <i>ICASSP 2025 - 2025 IEEE Interna-</i>	dividual differences: comparability and incompa-	1448
1396	<i>tional Conference on Acoustics, Speech and Signal</i>	rability across species. <i>Philosophical Transac-</i>	1449
1397	<i>Processing (ICASSP)</i> , pages 1–5.	<i>tions of the Royal Society B: Biological Sciences</i> ,	1450
1398	Jing Peng, Yucheng Wang, Yu Xi, Xu Li, Xizhuo Zhang,	373(1744):20170170.	1451
1399	and Kai Yu. 2024. A survey on speech large language	Emanuel A Schegloff. 1982. Discourse as an interac-	1452
1400	models. <i>arXiv preprint arXiv:2410.18908</i> .	tional achievement: Some uses of ‘uh huh’ and other	1453
1401	Sabid Bin Habib Pias, Alicia Freel, Ran Huang, Donald	things that come between sentences. <i>Analyzing dis-</i>	1454
1402	Williamson, Minjeong Kim, and Apu Kapadia. 2024.	<i>course: Text and talk</i> , 71(93).	1455
1403	Building trust through voice: How vocal tone impacts	Emanuel A Schegloff. 2000. Overlapping talk and the	1456
1404	user perception of attractiveness of voice assistants.	organization of turn-taking for conversation. <i>Lan-</i>	1457
1405	<i>arXiv preprint arXiv:2409.18941</i> .	<i>guage in society</i> , 29(1):1–63.	1458
1406	Karol J. Piczak. 2015a. ESC: Dataset for Environmental	Maureen Seyssel, Antony D’Avirro, Adina Williams,	1459
1407	Sound Classification. In <i>Proceedings of the 23rd</i>	and Emmanuel Dupoux. 2024. Emphassess: a	1460
1408	<i>Annual ACM Conference on Multimedia</i> , pages 1015–	prosodic benchmark on assessing emphasis transfer	1461
1409	1018. ACM Press.	in speech-to-speech models. In <i>Proceedings of the</i>	1462
1410	Karol J Piczak. 2015b. Esc: Dataset for environmental	<i>2024 Conference on Empirical Methods in Natural</i>	1463
1411	sound classification. In <i>Proceedings of the 23rd ACM</i>	<i>Language Processing</i> , pages 495–507.	1464
1412	<i>international conference on Multimedia</i> , pages 1015–	Xinyue Shen, Yixin Wu, Michael Backes, and Yang	1465
1413	1018.	Zhang. 2024. Voice jailbreak attacks against gpt-4o.	1466
1414	Vineel Pratap, Qiantong Xu, Anuroop Sriram, Gabriel	<i>arXiv preprint arXiv:2405.19103</i> .	1467
1415	Synnaeve, and Ronan Collobert. 2020. Mls: A large-	Jack Sidnell. 2007. Comparative studies in conversation	1468
1416	scale multilingual dataset for speech research. In	analysis. <i>Annu. Rev. Anthropol.</i> , 36(1):229–244.	1469
1417	<i>Proc. Interspeech 2020</i> , pages 2757–2761.	Sunayana Sitaram, Khyathi Raghavi Chandu, Sai Kr-	1470
1418	Alec Radford, Jong Wook Kim, Tao Xu, Greg Brock-	ishna Rallabandi, and Alan W Black. 2019. A survey	1471
1419	man, Christine McLeavey, and Ilya Sutskever. 2023.	of code-switched speech and language processing.	1472
1420	Robust speech recognition via large-scale weak su-	<i>arXiv preprint arXiv:1904.00784</i> .	1473
1421	perception. In <i>International conference on machine</i>	Juntao Tan, Liangwei Yang, Zuxin Liu, Zhiwei Liu,	1474
1422	<i>learning</i> , pages 28492–28518. PMLR.	Rithesh Murthy, Tulika Manoj Awalgaonkar, Jianguo	1475
1423	David Robinson, Marius Miron, Masato Hagiwara, and	Zhang, Weiran Yao, Ming Zhu, Shirley Kokane, and	1476
1424	Olivier Pietquin. 2025. NatureLM-audio: an audio-	1 others. 2025. Personabench: Evaluating ai models	1477
1425	language foundation model for bioacoustics . In <i>The</i>	on understanding personal information through ac-	1478
1426	<i>Thirteenth International Conference on Learning</i>	cessing (synthetic) private user data. <i>arXiv preprint</i>	1479
1427	<i>Representations</i> .	<i>arXiv:2502.20616</i> .	1480
1428	Jaechul Roh, Virat Shejwalkar, and Amir Houmansadr.	Zhaoxuan Tan, Zheyuan Liu, and Meng Jiang. 2024.	1481
1429	2025. Multilingual and multi-accent jailbreaking of	Personalized pieces: Efficient personalized large	1482
1430	audio llms. <i>arXiv preprint arXiv:2504.01094</i> .	language models through collaborative efforts. In <i>Pro-</i>	1483
		<i>ceedings of the 2024 Conference on Empirical Meth-</i>	1484
		<i>ods in Natural Language Processing</i> , pages 6459–	1485
		6475.	1486

1487	Changli Tang, Wenyi Yu, Guangzhi Sun, Xianzhao	Yingzhi Wang, Pooneh Mousavi, Artem Ploujnikov, and	1540
1488	Chen, Tian Tan, Wei Li, Lu Lu, Zejun MA, and Chao	Mirco Ravanelli. 2025d. What are they doing? joint	1541
1489	Zhang. 2024. SALMONN: Towards generic hearing	audio-speech co-reasoning . In <i>ICASSP 2025 - 2025</i>	1542
1490	abilities for large language models . In <i>The Twelfth</i>	<i>IEEE International Conference on Acoustics, Speech</i>	1543
1491	<i>International Conference on Learning Representa-</i>	<i>and Signal Processing (ICASSP)</i> , pages 1–5.	1544
1492	<i>tions</i> .		
1493	Gemini Team, Petko Georgiev, Ving Ian Lei, Ryan	Yubo Wang, Xueguang Ma, Ge Zhang, Yuansheng Ni,	1545
1494	Burnell, Libin Bai, Anmol Gulati, Garrett Tanzer,	Abhranil Chandra, Shiguang Guo, Weiming Ren,	1546
1495	Damien Vincent, Zhufeng Pan, Shibo Wang, and 1	Aaran Arulraj, Xuan He, Ziyang Jiang, Tianle Li, Max	1547
1496	others. 2024. Gemini 1.5: Unlocking multimodal	Ku, Kai Wang, Alex Zhuang, Rongqi Fan, Xiang	1548
1497	understanding across millions of tokens of context.	Yue, and Wenhui Chen. 2024. MMLU-pro: A more	1549
1498	<i>arXiv preprint arXiv:2403.05530</i> .	robust and challenging multi-task language under-	1550
1499		standing benchmark . In <i>The Thirty-eight Conference</i>	1551
1500	Joseph Turian, Jordie Shier, Humair Raj Khan, Bhiksha	<i>on Neural Information Processing Systems Datasets</i>	1552
1501	Raj, Björn W Schuller, Christian J Steinmetz, Colin	<i>and Benchmarks Track</i> .	1553
1502	Malloy, George Tzanetakis, Gissel Velarde, Kirk Mc-	Pete Warden. 2018. Speech commands: A dataset	1554
1503	Nally, and 1 others. 2022. Hear: Holistic evaluation	for limited-vocabulary speech recognition. <i>arXiv</i>	1555
1504	of audio representations. In <i>NeurIPS 2021 Compe-</i>	<i>preprint arXiv:1804.03209</i> .	1556
1505	<i>titions and Demonstrations Track</i> , pages 125–145.		
1506	PMLR.	Benno Weck, Ilaria Manco, Emmanouil Benetos, Elio	1557
1507	George Tzanetakis and Perry Cook. 2002. Musical	Quinton, György Fazekas, and Dmitry Bogdanov.	1558
1508	genre classification of audio signals. <i>IEEE Trans-</i>	2024. Muchomusic: Evaluating music understand-	1559
1509	<i>actions on speech and audio processing</i> , 10(5):293–	ing in multimodal audio-language models. In <i>Pro-</i>	1560
1510	302.	<i>ceedings of the 25th International Society for Music</i>	1561
1511	Jörgen Valk and Tanel Alumäe. 2021. Voxlingua107:	<i>Information Retrieval Conference (ISMIR)</i> .	1562
1512	a dataset for spoken language recognition. In <i>2021</i>		
1513	<i>IEEE Spoken Language Technology Workshop (SLT)</i> ,	Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten	1563
1514	pages 652–658. IEEE.	Bosma, Fei Xia, Ed Chi, Quoc V Le, Denny Zhou,	1564
1515	Christophe Veaux, Junichi Yamagishi, Kirsten MacDon-	and 1 others. 2022. Chain-of-thought prompting elic-	1565
1516	ald, and 1 others. 2017. Cstr vctk corpus: English	its reasoning in large language models. <i>Advances</i>	1566
1517	multi-speaker corpus for cstr voice cloning toolkit.	<i>in neural information processing systems</i> , 35:24824–	1567
1518	<i>University of Edinburgh. The Centre for Speech Tech-</i>	24837.	1568
1519	<i>nology Research (CSTR)</i> , 6:15.	Victor Junqiu Wei, Weicheng Wang, Di Jiang, Yuanfeng	1569
1520	Denny Vrandečić and Markus Kröttsch. 2014. Wiki-	Song, and Lu Wang. 2024. Asr-ec benchmark: Eval-	1570
1521	data: a free collaborative knowledgebase. <i>Communi-</i>	uating large language models on chinese asr error	1571
1522	<i>cations of the ACM</i> , 57(10):78–85.	correction. <i>arXiv preprint arXiv:2412.03075</i> .	1572
1523	Bin Wang, Xunlong Zou, Geyu Lin, Shuo Sun, Zhuohan	Candace West. 1979. Against our will: Male interrup-	1573
1524	Liu, Wenyu Zhang, Zhengyuan Liu, AiTi Aw, and	tions of females in cross-sex conversation. <i>Annals of</i>	1574
1525	Nancy F Chen. 2025a. Audiobench: A universal	<i>the New York Academy of Sciences</i> .	1575
1526	benchmark for audio large language models. <i>NAACL</i> .		
1527	Bin Wang, Xunlong Zou, Shuo Sun, Wenyu Zhang,	Haibin Wu, Xuanjun Chen, Yi-Cheng Lin, Kai-wei	1576
1528	Yingxu He, Zhuohan Liu, Chengwei Wei, Nancy F	Chang, Ho-Lam Chung, Alexander H Liu, and Hung-	1577
1529	Chen, and AiTi Aw. 2025b. Advancing singlish un-	yi Lee. 2024a. Towards audio language modeling—an	1578
1530	derstanding: Bridging the gap with datasets and mul-	overview. <i>arXiv preprint arXiv:2402.13236</i> .	1579
1531	<i>timodal models</i> . <i>arXiv preprint arXiv:2501.01034</i> .		
1532	Changhan Wang, Anne Wu, Jiatao Gu, and Juan Pino.	Junkai Wu, Xulin Fan, Bo-Ru Lu, Xilin Jiang, Nima	1580
1533	2021. Covost 2 and massively multilingual speech	Mesgarani, Mark Hasegawa-Johnson, and Mari Os-	1581
1534	translation . In <i>Interspeech 2021</i> , pages 2247–2251.	tendorf. 2024b. Just asr+ llm? a study on speech	1582
1535	Siyin Wang, Wenyi Yu, Xianzhao Chen, Xiaohai Tian,	large language models’ ability to identify and un-	1583
1536	Jun Zhang, Yu Tsao, Junichi Yamagishi, Yuxuan	derstand speaker in spoken dialogue. In <i>2024 IEEE</i>	1584
1537	Wang, and Chao Zhang. 2025c. Qualispeech: A	<i>Spoken Language Technology Workshop (SLT)</i> , pages	1585
1538	speech quality assessment dataset with natural lan-	1137–1143. IEEE.	1586
1539	guage reasoning and descriptions. <i>arXiv preprint</i>	Erjia Xiao, Hao Cheng, Jing Shao, Jinhao Duan, Kaidi	1587
	<i>arXiv:2503.20290</i> .	Xu, Le Yang, Jindong Gu, and Renjing Xu. 2025.	1588
		Tune in, act up: Exploring the impact of audio	1589
		modality-specific edits on large audio language mod-	1590
		els in jailbreak. <i>arXiv preprint arXiv:2501.13772</i> .	1591
		Liumeng Xue, Ziya Zhou, Jiahao Pan, Zixuan Li,	1592
		Shuai Fan, Yinghao Ma, Sitong Cheng, Dongchao	1593
		Yang, Haohan Guo, Yujia Xiao, and 1 others. 2025.	1594
		Audio-flan: A preliminary release. <i>arXiv preprint</i>	1595
		<i>arXiv:2502.16584</i> .	1596

1597	Ruiqi Yan, Xiquan Li, Wenxi Chen, Zhikang Niu,	Shukang Yin, Chaoyou Fu, Sirui Zhao, Ke Li, Xing	1655
1598	Chen Yang, Ziyang Ma, Kai Yu, and Xie Chen.	Sun, Tong Xu, and Enhong Chen. 2024. A survey on	1656
1599	2025. Uro-bench: A comprehensive benchmark for	multimodal large language models. <i>National Science</i>	1657
1600	end-to-end spoken dialogue models. <i>arXiv preprint</i>	<i>Review</i> , 11(12).	1658
1601	<i>arXiv:2502.17810</i> .		
1602	Chih-Kai Yang, Yu-Kuan Fu, Chen-An Li, Yi-Cheng	Ruibin Yuan, Yinghao Ma, Yizhi Li, Ge Zhang, Xingran	1659
1603	Lin, Yu-Xiang Lin, Wei-Chih Chen, Ho Lam Chung,	Chen, Hanzhi Yin, Yiqi Liu, Jiawen Huang, Zeyue	1660
1604	Chun-Yi Kuan, Wei-Ping Huang, Ke-Han Lu, and 1	Tian, Binyue Deng, and 1 others. 2023. Marble:	1661
1605	others. 2024a. Building a taiwanese mandarin spo-	Music audio representation benchmark for universal	1662
1606	ken language model: A first attempt. <i>arXiv preprint</i>	evaluation. <i>Advances in Neural Information Process-</i>	1663
1607	<i>arXiv:2411.07111</i> .	<i>ing Systems</i> , 36:39626–39647.	1664
1608	Chih-Kai Yang, Neo Ho, Yen-Ting Piao, and Hung-yi	Yongyi Zang, Sean O’Brien, Taylor Berg-Kirkpatrick,	1665
1609	Lee. 2025a. Sakura: On the multi-hop reasoning of	Julian McAuley, and Zachary Novack. 2025. Are	1666
1610	large audio-language models based on speech and	you really listening? boosting perceptual aware-	1667
1611	audio information. <i>Interspeech 2025</i> .	ness in music-qa benchmarks. <i>arXiv preprint</i>	1668
		<i>arXiv:2504.00369</i> .	1669
1612	Chih-Kai Yang, Kuan-Po Huang, Ke-Han Lu, Chun-	Heiga Zen, Viet Dang, Rob Clark, Yu Zhang, Ron J	1670
1613	Yi Kuan, Chi-Yuan Hsiao, and Hung-Yi Lee.	Weiss, Ye Jia, Zhifeng Chen, and Yonghui Wu. 2019.	1671
1614	2024b. Investigating zero-shot generalizability on	Libritts: A corpus derived from librispeech for text-	1672
1615	mandarin-english code-switched asr and speech-to-	to-speech. In <i>Proc. Interspeech 2019</i> , pages 1526–	1673
1616	text translation of recent foundation models with self-	1530.	1674
1617	supervision and weak supervision. In <i>2024 IEEE</i>		
1618	<i>International Conference on Acoustics, Speech, and</i>	Zhehao Zhang, Ryan A Rossi, Branislav Kveton, Yijia	1675
1619	<i>Signal Processing Workshops (ICASSPW)</i> , pages 540–	Shao, Diyi Yang, Hamed Zamani, Franck Dernon-	1676
1620	544.	court, Joe Barrow, Tong Yu, Sungchul Kim, and 1	1677
		others. 2024. Personalization of large language mod-	1678
1621	Hao Yang, Lizhen Qu, Ehsan Shareghi, and Gholamreza	els: A survey. <i>arXiv preprint arXiv:2411.00027</i> .	1679
1622	Haffari. 2025b. Audio is the achilles’ heel: Red team-		
1623	ing audio large multimodal models. In <i>Proceedings</i>	Mengjie Zhao, Zhi Zhong, Zhuoyuan Mao, Shiqi	1680
1624	<i>of the 2025 Conference of the Nations of the Americas</i>	Yang, Wei-Hsiang Liao, Shusuke Takahashi, Hiromi	1681
1625	<i>Chapter of the Association for Computational Linguis-</i>	Wakaki, and Yuki Mitsufuji. 2024. Openmu: Your	1682
1626	<i>tistics: Human Language Technologies (Volume 1: Long</i>	swiss army knife for music understanding. <i>arXiv</i>	1683
1627	<i>Papers)</i> , pages 9292–9306, Albuquerque, New	<i>preprint arXiv:2410.15573</i> .	1684
1628	Mexico. Association for Computational Linguistics.		
1629	Qian Yang, Jin Xu, Wenrui Liu, Yunfei Chu, Ziyue	Wayne Xin Zhao, Kun Zhou, Junyi Li, Tianyi Tang,	1685
1630	Jiang, Xiaohuan Zhou, Yichong Leng, Yuanjun Lv,	Xiaolei Wang, Yupeng Hou, Yingqian Min, Beichen	1686
1631	Zhou Zhao, Chang Zhou, and Jingren Zhou. 2024c.	Zhang, Junjie Zhang, Zican Dong, and 1 others. 2023.	1687
1632	AIR-bench: Benchmarking large audio-language	A survey of large language models. <i>arXiv preprint</i>	1688
1633	models via generative comprehension. In <i>Proceed-</i>	<i>arXiv:2303.18223</i> .	1689
1634	<i>ings of the 62nd Annual Meeting of the Association</i>	Jeffrey Zhou, Tianjian Lu, Swaroop Mishra, Sid-	1690
1635	<i>for Computational Linguistics (Volume 1: Long Pa-</i>	dhartha Brahma, Sujoy Basu, Yi Luan, Denny Zhou,	1691
1636	<i>pers)</i> , pages 1979–1998, Bangkok, Thailand. Associ-	and Le Hou. 2023a. Instruction-following evalu-	1692
1637	ation for Computational Linguistics.	ation for large language models. <i>arXiv preprint</i>	1693
		<i>arXiv:2311.07911</i> .	1694
1638	Shu-wen Yang, Po-Han Chi, Yung-Sung Chuang,	Kun Zhou, Yutao Zhu, Zhipeng Chen, Wentong Chen,	1695
1639	Cheng-I Jeff Lai, Kushal Lakhotia, Yist Y Lin,	Wayne Xin Zhao, Xu Chen, Yankai Lin, Ji-Rong	1696
1640	Andy T Liu, Jiatong Shi, Xuankai Chang, Guan-Ting	Wen, and Jiawei Han. 2023b. Don’t make your llm	1697
1641	Lin, and 1 others. 2021. Superb: Speech processing	an evaluation benchmark cheater. <i>arXiv preprint</i>	1698
1642	universal performance benchmark. In <i>Proc. Inter-</i>	<i>arXiv:2311.01964</i> .	1699
1643	<i>speech 2021</i> , pages 1194–1198.		
1644	Sohee Yang, Elena Gribovskaya, Nora Kassner, Mor	Andy Zou, Zifan Wang, Nicholas Carlini, Milad Nasr,	1700
1645	Geva, and Sebastian Riedel. 2024d. Do large lan-	J Zico Kolter, and Matt Fredrikson. 2023. Univer-	1701
1646	guage models latently perform multi-hop reasoning?	saral and transferable adversarial attacks on aligned	1702
1647	In <i>Proceedings of the 62nd Annual Meeting of the</i>	language models. <i>arXiv preprint arXiv:2307.15043</i> .	1703
1648	<i>Association for Computational Linguistics (Volume</i>		
1649	<i>1: Long Papers)</i> , pages 10210–10229.		
1650	Sohee Yang, Nora Kassner, Elena Gribovskaya, Se-	A Detailed Categorization of the	1704
1651	bastian Riedel, and Mor Geva. 2024e. Do large	Surveyed Papers	1705
1652	language models perform latent multi-hop reason-		
1653	ing without exploiting shortcuts? <i>arXiv preprint</i>	The complete categorization of the surveyed papers,	1706
1654	<i>arXiv:2411.16679</i> .	based on the proposed taxonomy (§2), is presented	1707
		in Figure 3. Please note that widely used corpora	1708

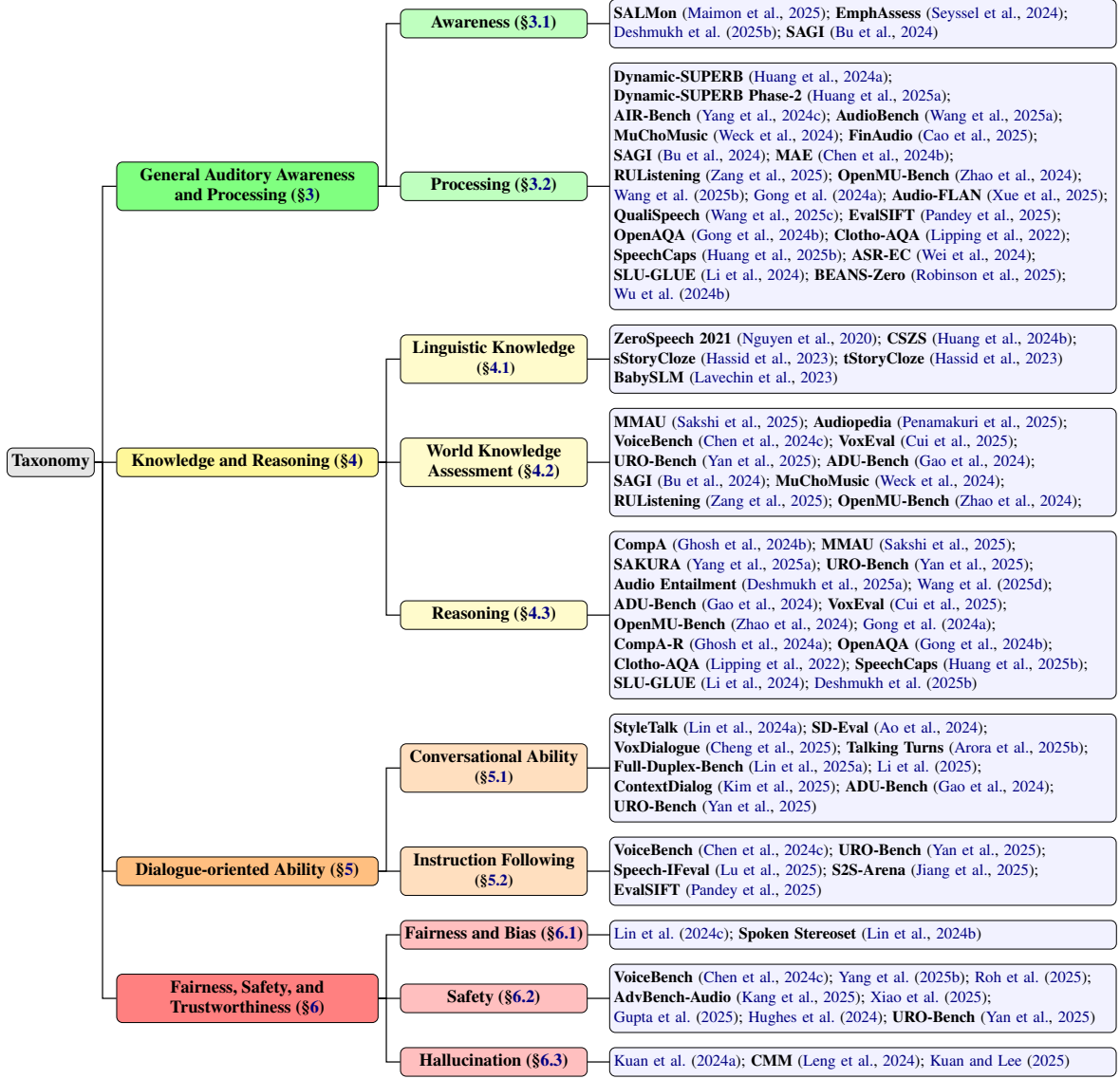


Figure 3: The complete categorization of the surveyed papers based on the proposed taxonomy.

for fundamental auditory processing tasks, such as speech recognition and audio captioning, are excluded from this categorization due to the extremely large number of such resources. Including them would make the figure overly detailed and cumbersome. For reference, we provide examples of these fundamental tasks and their corresponding resources in Appendix B.

From Figure 3, it is evident that the current focus of LALM evaluations predominantly centers on auditory processing tasks (§3.2), underscoring their importance to the research community. While these tasks are undeniably valuable, they should not be seen as the sole consideration when evaluating models for real-world applications. A more diverse and comprehensive evaluation scope is crucial to ensure a fuller understanding of their potential and

shortcomings.

B Examples of General Auditory Processing Tasks

Table 1 lists representative auditory processing tasks and their associated resources. As foundational components of auditory processing, these tasks are well-suited for adaptation in LALM evaluation, as discussed in (§3.2).

C Dynamics in Full-Duplex Dialogues

In this section, we briefly introduce the dynamics discussed in (§5.1.2). **Turn-taking** (Sacks et al., 1974) is a fundamental aspect of conversational organization, where speakers alternate turns to speak, ensuring only one person talks at a time. This process is complex, involving various behaviors that

Auditory Tasks	Common Datasets
Audio Tasks	
Audio Captioning	AudioCaps (Kim et al., 2019) Clotho (Drossos et al., 2020)
Audio Classification	ESC-50 (Piczak, 2015b) AudioSet (Gemmeke et al., 2017)
Vocal Sound Classification	VocalSound (Gong et al., 2022)
Speech Tasks	
Automatic Speech Recognition	LibriSpeech (Panayotov et al., 2015b) AISHELL-1 (Bu et al., 2017) Common Voice (Ardila et al., 2020)
Speaker Identification	VoxCeleb2 (Chung et al., 2018) CN-Celeb (Fan et al., 2020)
Text-to-Speech	LJSpeech (Ito and Johnson, 2017) VCTK (Veaux et al., 2017) LibriTTS (Zen et al., 2019)
Speech Emotion Recognition	IEMOCAP (Busso et al., 2008) CREMA-D (Cao et al., 2014)
Language Identification	VoxLingua107 (Valk and Alumäe, 2021) FLEURS (Conneau et al., 2023)
Speech Translation	CoVoST 2 (Wang et al., 2021) MuST-C (Di Gangi et al., 2019)
Speech Diarization	LibriMix (Cosentino et al., 2020)
Keyword Spotting	Speech Command (Warden, 2018)
Music Tasks	
Music Captioning Text-to-Music	MusicCaps (Agostinelli et al., 2023) Song Descriptor Dataset (Manco et al., 2023) MidiCaps (Melechovsky et al., 2024)
Music Transcription	MAESTRO (Hawthorne et al., 2019)
Instrument Classification	NSynth (Engel et al., 2017)
Genre Classification	FMA (Defferrard et al., 2017) GTZAN (Tzanetakis and Cook, 2002)

Table 1: Commonly used datasets for various auditory tasks. This overview covers key tasks in audio, speech, and music processing and the datasets that are widely adopted in academic and industrial research.

help facilitate smooth transitions between speakers. For example, speakers often signal the end of their turn through clear cues, allowing the listener to recognize when they are yielding the floor (Duncan, 1972; Duncan and Fiske, 2015). Furthermore, turn-taking conventions may be shaped by cultural factors (Sidnell, 2007), which influence how and when speakers take their turns due to linguistic and social differences. Understanding and modeling these behaviors are essential steps toward achieving natural and effective communication in both human-human and human-AI interactions.

Backchanneling involves the listener’s use of phatic expressions that signal active listening and attentiveness to the speaker (Fujie et al., 2005). These brief verbal cues, such as “yeah,” “I see,” or “uh-huh,” along with non-verbal cues like nodding, serve as feedback, demonstrating sympathy, agreement, or understanding. By offering such responses, listeners help maintain the flow of conversation without interrupting the speaker. This

behavior not only fosters a sense of connection but also enhances the speaker’s feeling of being heard and understood, contributing to a more interactive and supportive dialogue. As such, backchanneling plays a crucial role in sustaining conversation dynamics and promoting positive communicative exchanges.

Speaker overlap refers to the simultaneous speech of multiple speakers, while **speaker interruption** occurs when one speaker interjects during another’s turn, which breaks the turn-taking principles (Gravano and Hirschberg, 2012). These phenomena are complex: they can be competitive, reflecting hostility or dominance (West, 1979; Orcutt and Harvey, 1985), or they can be neutral or supportive, helping to maintain and coordinate the flow of dialogue (Goldberg, 1990; Jefferson, 1986; Gervits and Scheutz, 2018). Despite their varying forms, both overlap and interruption are natural components of human conversation.

D Input/Output Modalities of the Surveyed Works

Our proposed taxonomy (§2) is organized by the evaluation objectives of the surveyed works rather than by the modalities they cover. Nevertheless, modality information is essential for researchers seeking benchmarks suited to models specialized in particular modalities. Thus, we provide the input/output modality details in Tables 2, 3, 4, and 5, corresponding to the categories of General Auditory Awareness and Processing (§3), Knowledge and Reasoning (§4), Dialogue-oriented Ability (§5), and Fairness, Safety, and Trustworthiness (§6), respectively. These tables are compiled based on the original papers of the surveyed benchmarks.

Please note that due to unique evaluation designs, some benchmarks do not produce explicit “outputs” but instead rely on input likelihood comparisons or similarity measures with specific instances. This absence of outputs is clearly indicated in the tables.

E Information of AI Assistance in Revision

We acknowledge the assistance of GPT-4.1-mini in refining the paper and improving its clarity.

General Auditory Awareness and Processing								
Benchmark	Input Modalities				Output Modalities			
	Text	Audio	Speech	Music	Text	Audio	Speech	Music
SALMon (Maimon et al., 2025)		✓	✓		Likelihood-based evaluation. No output modality.			
Wu et al. (2024b)	✓		✓		✓			
ADU-Bench (Gao et al., 2024)			✓		✓		✓	
EmphAssess (Seyssel et al., 2024)			✓				✓	
Deshmukh et al. (2025b)	✓	✓	✓		✓			
Dynamic-SUPERB (Huang et al., 2024a)	✓	✓	✓	✓	✓			
Dynamic-SUPERB Phase-2 (Huang et al., 2025a)	✓	✓	✓	✓	✓			
AIR-Bench (Yang et al., 2024c)	✓	✓	✓	✓	✓			
AudioBench (Wang et al., 2025a)	✓	✓	✓		✓			
MuChoMusic (Weck et al., 2024)	✓			✓	✓			
FinAudio (Cao et al., 2025)	✓		✓		✓			
SAGI (Bu et al., 2024)	✓	✓	✓	✓	✓			
MAE (Chen et al., 2024b)	✓	✓	✓		✓			
RUListing (Zang et al., 2025)	✓			✓	✓			
OpenMU-Bench (Zhao et al., 2024)	✓			✓	✓			
Wang et al. (2025b)	✓		✓		✓			
Gong et al. (2024a)	✓	✓	✓	✓	✓			
Audio-FLAN (Xue et al., 2025)	✓	✓	✓	✓	✓	✓	✓	✓
QualiSpeech (Wang et al., 2025c)	✓		✓		✓			
EvalSIFT (Pandey et al., 2025)	✓		✓		✓		✓	
OpenAQA (Gong et al., 2024b)	✓	✓			✓			
Clotho-AQA (Lipping et al., 2022)	✓	✓			✓			
SpeechCaps (Huang et al., 2025b)	✓		✓		✓			
ASR-EC (Wei et al., 2024)	✓		✓		✓			
SLU-GLUE (Li et al., 2024)	✓		✓		✓			
BEANS-Zero (Robinson et al., 2025)	✓			✓	✓			

Table 2: Input and output modalities of benchmarks in the **General Auditory Awareness and Processing** category shown in Figure 3.

Knowledge and Reasoning								
Benchmark	Input Modalities				Output Modalities			
	Text	Audio	Speech	Music	Text	Audio	Speech	Music
ZeroSpeech 2021 (Nguyen et al., 2020)			✓		Likelihood-based evaluation. No output modality.			
CSZS (Huang et al., 2024b)			✓		Likelihood-based evaluation. No output modality.			
sStoryCloze (Hassid et al., 2023)			✓		Likelihood-based evaluation. No output modality.			
tStoryCloze (Hassid et al., 2023)			✓		Likelihood-based evaluation. No output modality.			
BabySLM (Lavechin et al., 2023)	✓		✓		Likelihood-based evaluation. No output modality.			
CompA (Ghosh et al., 2024b)	✓	✓			Similarity-based evaluation on text and audio inputs.			
MMAU (Sakshi et al., 2025)	✓	✓	✓	✓	✓			
Audiopedia (Penamakuri et al., 2025)	✓		✓		✓			
VoiceBench (Chen et al., 2024c)	✓		✓		✓			
VoxEval (Cui et al., 2025)			✓				✓	
SAKURA (Yang et al., 2025a)	✓	✓	✓		✓			
URO-Bench (Yan et al., 2025)	✓		✓		✓		✓	
Audio Entailment (Deshmukh et al., 2025a)	✓	✓			✓			
ADU-Bench (Gao et al., 2024)			✓		✓		✓	
SAGI (Bu et al., 2024)	✓	✓	✓	✓	✓			
MuChoMusic (Weck et al., 2024)	✓			✓	✓			
RUListing (Zang et al., 2025)	✓			✓	✓			
OpenMU-Bench (Zhao et al., 2024)	✓			✓	✓			
Gong et al. (2024a)	✓	✓	✓	✓	✓			
CompA-R (Ghosh et al., 2024a)	✓	✓			✓			
OpenAQA (Gong et al., 2024b)	✓	✓			✓			
Clotho-AQA (Lipping et al., 2022)	✓	✓			✓			
SLU-GLUE (Li et al., 2024)	✓		✓		✓			
SpeechCaps (Huang et al., 2025b)	✓		✓		✓			
Wang et al. (2025d)	✓	✓	✓		✓			
Deshmukh et al. (2025b)	✓	✓	✓		✓			

Table 3: Input and output modalities of benchmarks in the **Knowledge and Reasoning** category shown in Figure 3.

Dialogue-oriented Ability								
Benchmark	Input Modalities				Output Modalities			
	Text	Audio	Speech	Music	Text	Audio	Speech	Music
StyleTalk (Lin et al., 2024a)	✓		✓				✓	
SD-Eval (Ao et al., 2024)	✓	✓	✓		✓			
VoxDialogue (Cheng et al., 2025)		✓	✓	✓	✓			
Talking Turns (Arora et al., 2025b)			✓				✓	
Full-Duplex-Bench (Lin et al., 2025a)			✓				✓	
Li et al. (2025)			✓		✓			
ContextDialog (Kim et al., 2025)			✓		✓		✓	
ADU-Bench (Gao et al., 2024)			✓		✓		✓	
VoiceBench (Chen et al., 2024c)	✓		✓		✓			
URO-Bench (Yan et al., 2025)	✓		✓		✓		✓	
Speech-IFeval (Lu et al., 2025)	✓		✓		✓			
S2S-Arena (Jiang et al., 2025)			✓				✓	
EvalSIFT (Pandey et al., 2025)	✓		✓		✓		✓	

Table 4: Input and output modalities of benchmarks in the **Dialogue-oriented Ability** category shown in Figure 3.

Fairness, Safety, and Trustworthiness								
Benchmark	Input Modalities				Output Modalities			
	Text	Audio	Speech	Music	Text	Audio	Speech	Music
Lin et al. (2024c)	✓		✓		✓			
Spoken Stereoset (Lin et al., 2024b)	✓		✓		✓			
VoiceBench (Chen et al., 2024c)	✓		✓		✓			
Yang et al. (2025b)	✓	✓	✓		✓			
Roh et al. (2025)	✓		✓		✓			
AdvBench-Audio (Kang et al., 2025)	✓		✓		✓			
Xiao et al. (2025)	✓		✓		✓			
Gupta et al. (2025)	✓		✓		✓			
Hughes et al. (2024)	✓		✓		✓			
URO-Bench (Yan et al., 2025)	✓		✓		✓		✓	
Kuan et al. (2024a)	✓	✓	✓		✓			
CMM (Leng et al., 2024)	✓	✓			✓			
Kuan and Lee (2025)	✓	✓			✓			

Table 5: Input and output modalities of benchmarks in the **Fairness, Safety, and Trustworthiness** category shown in Figure 3.