
Same Words, Different Judgments: How Preferences Vary Across Modalities

Aaron Broukhim

Department of Computer Science and Engineering
University of California San Diego
La Jolla, CA 92093
aabroukh@ucsd.edu

Nadir Weibel

Department of Computer Science and Engineering
University of California San Diego
La Jolla, CA 92093
weibel@ucsd.edu

Eshin Jolly

Department of Psychology
University of California San Diego
La Jolla, CA 92093
e3jolly@ucsd.edu

Abstract

Preference-based reinforcement learning (PbRL) is the dominant framework for aligning AI systems to human preferences. However, evaluation protocols for such data were designed for text and have not been validated for speech. We present the first ICC-based, controlled cross-modal study of human and synthetic preference annotations, comparing text and audio evaluations of identical semantic content across 100 prompts. We show that achieving *good* agreement within either modality ($\text{ICC}(2,k) \approx .80$) requires ~ 9 raters. At the same time, modalities show marked differences in how people report preferences: audio raters exhibit narrower decision thresholds, reduced length bias, and more user-oriented evaluation criteria, with near-chance cross-modality agreement. We demonstrate that synthetic ratings can be used to effectively predict inter-rater agreement, thus serving as an early signal for stimulus selection and proxy for human annotations. Together, these findings argue that evaluation protocols for audio preference data require modality-specific design rather than direct adaptation from text.¹

1 Introduction

Speech audio models have seen rapid development and deployment in recent years, yet research on aligning them to human preferences remains sparse [3]. Preference-based reinforcement learning (PbRL) offers a path forward, particularly for qualities like naturalness [14, 28], emotional expressiveness [12], and conversational appropriateness [7], where vocal delivery may shape perception alongside semantic content.

Current preference alignment research centers on text, following a standard pipeline: collect pairwise preferences, train a reward model, and optimize via RL [4]. This approach has revealed biases such as annotators’ tendency to prefer longer responses [22], while largely ignoring additional factors like attention allocation and reading order. Audio by contrast, delivers information at a controlled, recordable pace but requires sequential presentation, making *order* an explicit source of influence on annotations.

¹Code, data, and the audio corpus are available at: <https://huggingface.co/datasets/NeurIPS-Anon-2784/modality-prefs-data>

In this work, we empirically investigate: (1) How reliable are audio preferences? (2) How does modality affect preference ratings? (3) Can synthetic ratings substitute or augment human ratings? Our contribution is threefold: (a) an evaluation protocol for cross-modal preference elicitation (continuous VAS ratings, counterbalanced sequential audio presentation, attention checks adapted for audio); (b) an empirical characterization of how reliability scales with rater count and how preference judgments shift across modalities; and (c) a released corpus of 3,113 TTS-converted conversations from PRISM to support future audio preference research. Because most existing PbRL work for speech relies on Text-To-Speech (TTS)-generated audio from existing text datasets, *assumed* to be appropriate for speech use cases [3], we focused our study on TTS-rendered speech of text originally written for reading, and address open questions in Section 6.

We find audio preference ratings as reliable as text but behaviorally distinct: they introduce recency effects while showing reduced length-bias susceptibility. We provide the first ICC-based characterization of preference-annotation reliability in either modality, showing that one annotator yields poor reliability, three moderate, and nine good agreement. Cross-modality agreement is near chance, with preference shifts varying by prompt rather than uniformly, suggesting TTS-converted text data may not be appropriate for audio preference pipelines. Synthetic ratings can replace or triage ambiguous response pairs, reducing annotation costs without sacrificing label quality.

2 Background

2.1 Preference-Based Reinforcement Learning

Preference-Based Reinforcement Learning is a framework in which a reward function is learned from human preferences rather than being explicitly specified. Given pairs (or sets) of trajectory segments or outputs, annotators indicate which they prefer, and these comparisons are used to train a reward model [6]. This reward model then serves as the optimization objective for a reinforcement learning agent. PbRL has been prominently applied to align Large Language Models with human values through Reinforcement Learning from Human Feedback (RLHF), where preference data is collected by having humans or LLMs compare candidate responses to the same query [4].

Standard PbRL pipelines typically rely on binary pairwise comparisons modeled with Bradley–Terry reward formulations [6]. While effective, binary labels discard information about preference strength and are susceptible to well-known order and position effects [2], which are difficult to diagnose without additional modeling assumptions or counterbalancing.

Recent work suggests richer scalar feedback yields more informative reward signals than binary comparisons [27], motivating continuous rating paradigms.

Known biases include length bias [22] (longer responses win regardless of content), position bias [26] (first-presented responses are favored), and reward overoptimization [11]. In audio, these biases are underexplored [3], and many works convert text datasets to audio via TTS [10], implicitly assuming text-suitable data transfers to speech—an open question.

2.2 Modality’s Effect on Judgments

Existing work on annotation biases in the speech literature has largely focused on quality assessment, cataloguing biases and proposing mitigations [30]. Explicit cross-modality comparisons of preference judgments remain rare. Most similar to this work, Cho et al. [5] demonstrated that preferences differ by modality and used this finding to build speech-adapted models, but their analysis did not extend to reliability scaling, order effects, decision thresholds, or synthetic-human alignment—dimensions we address here.

More broadly, speech processing is strictly linear and imposes higher cognitive load than reading [5], and users tend to anthropomorphize audio interfaces more readily than text-based ones [9], suggesting that evaluation criteria may shift across modalities.

2.3 Reliability and Agreement in Preference Annotations

Prior work on preference annotation reliability has generally reported only single-rater pairwise agreement—e.g., 73–77% for text summarization [24], 63% for helpful and harmless annotations [1],

and 60% for music [8]. While intuitive, these percentages are not directly comparable across studies and, critically, do not address how agreement scales with the number of raters. Intraclass correlation coefficients and Krippendorff’s α offer richer characterizations by partitioning variance across raters and stimuli, and by accommodating different measurement scales, respectively [17]. To our knowledge, no prior work has used these metrics to characterize the number of raters needed to reach specific reliability thresholds in preference annotation for either text or audio.

2.4 The PRISM Dataset

PRISM [16] is a preference dataset of 8,011 live conversations between 1,500 diverse participants and 21 LLMs. Participants rate each response on a 1–100 visual analog scale (VAS) with hidden numeric values to avoid anchoring. Crucially, the person who initiates each conversation is also the one who rates the responses, coupling the evaluative and interactive contexts. We use PRISM’s unguided conversations as our source material, adopting its VAS methodology and leveraging its original interactive ratings as our baseline.

3 Methods

We converted a subset of PRISM [16] to spoken audio using Kokoro [13] and collected matched text-to-text and audio-to-audio comparisons. Audio-condition participants also rated audio quality to control for fidelity confounds. This study was approved as minimal-risk research by the authors’ Institutional Review Board (IRB).

3.1 Converting Text to Audio

We sub-selected "unguided" conversations to maintain a more objective dataset, since other conversations were intended to be controversial. From the 3,113 unguided conversations, we randomly selected 100 interactions. To ensure variability in response lengths, 25% (25/100) of the interactions were drawn from the top 10th percentile of absolute character length differences between responses (i.e., a minimum difference of 449 characters). The rest (75/100) were drawn from the lower 90th percentile. This yielded 200 audio clips (2 model responses $\times$ 100 interactions). Audio clips were also screened so that only innocuous content was presented to participants, removing clips that could be construed as offensive, discriminatory, or emotionally distressing. Together, this allowed us to experimentally isolate the effect of modality and order on preference ratings.

Each interaction was passed through the current state-of-the-art open-source TTS model, Kokoro [13], as benchmarked by TTS-Arena [25], to generate audio samples. We used Kokoro’s default parameters, with the voice set to `af_heart` selected for its neutral, clear delivery and kept this consistent across all samples. We note that PRISM responses were originally written by LLMs for text presentation; spoken delivery of text-optimized content may differ from speech written for spoken delivery. Our findings therefore characterize the common pipeline of TTS-converting text-collected data, not the upper bound of audio-native preference annotation (see Section 6). All audio snippets were manually reviewed for artifacts or mispronunciations that could negatively impact quality to focus on semantic content. Low audio quality samples were replaced, and replacements were re-evaluated following the same selection criteria. We provide [anonymized documentation], publicly release [anonymized audio corpus] of all 3,113 conversations, and [anonymized code repository] for data collection and analysis to encourage audio preference dataset collection.

3.2 Data Collection

3.2.1 Demographic Task

Prior to preference solicitation, participants completed a demographic questionnaire including: age, gender, race, ethnicity, birthplace, education level, employment status, country of residence, primary language, English proficiency (for non-native speakers), other spoken languages, and current timezone. Participants with no English proficiency and those under 18 were excluded. Participant demographics can be found in Appendix Table 6.

3.2.2 Preference Tasks

For each interaction, a history of the conversation up to that point in text was shown to the user. The interface can be seen in Appendix Figures 3 and 4. Participants were then asked to rate the responses from 1 (Terrible) to 100 (Perfect) using a slider that does not present the numerical value to avoid anchoring, as done in the original PRISM paper [16]. Each participant rated 20 interactions presented in a balanced block random order with an attention check at the beginning (21 total). The attention check paired a coherent response with a clearly nonsensical one; participants who preferred the nonsensical response were excluded.

Participants were randomly assigned to either *text* or *audio* response modality to prevent cross-condition contamination. Those in the audio condition also rated audio quality using the same 1-100 scale. Audio quality ratings were collected to monitor for fidelity-related confounds in the response quality judgments. As expected, perceived audio quality was a positive predictor of response ratings within stimuli (see Appendix Section F), but controlling for it did not change the order, length, or trial-number effects reported in the main text. We therefore do not include audio-quality ratings as a co-variate in the main models.

Binary preference data, with the option to tie, were presented to respondents at the end. An optional text box asked raters why they preferred one response over the other for qualitative analysis. We presented both text and audio responses sequentially to explicitly compare order effects across modality.

Audio responses required full playback on the first listen—participants could not skip ahead or scrub through the timeline. After the initial playback, participants could replay the audio as many times as needed, ensuring parity with the text condition where responses could be re-read at will.

Beyond explicit ratings, we collected additional behavioral indicators (decision time, audio deliberation time, replay counts, intermediate rating changes including reversals after listening to both clips), session-level metadata (duration, completion, dropout), and trial order metadata (which response appeared first, trial position 1–20). We also recorded stimulus properties: history character count and turns, response character length, and absolute character-length differences between paired responses (character length is linearly correlated with audio length, so we treat these as interchangeable for analysis purposes).

3.2.3 Synthetic Ratings

We collected synthetic ratings (10 per prompt per modality, 1–100 scale) using GPT-4o for text [15] and GPT-4o-Audio-Preview for audio with default API parameters. We chose GPT-4o-Audio-Preview as an end-to-end audio model rather than an STT+LLM+TTS pipeline that would only evaluate transcription/synthesis accuracy. These synthetic scores let us test whether AI signals can augment or replace human annotations for Reinforcement Learning from AI Feedback (RLAIF) applications.

3.3 Participants and Exclusions

We recruited 106 participants (53 per modality) from Prolific’s [20] US population at \$12/hour. 7/53 (13.2%) audio and 5/53 (9.4%) text participants failed the attention check; of the remainder, 4 text participants did not finish the survey but we retained their completed trials, yielding 94 participants and 8–10 ratings per audio clip. Demographics are reported in Appendix Table 6.

4 Results

4.1 Cross-Modal Rating Characteristics

Decision Threshold. Raters exhibited larger average rating differences when committing to a winner (not a tie) when presented as text ($M = 41.7$, 95% CI [39.5, 43.9]; $Mdn = 34.0$) compared to audio ($M = 27.9$, 95% CI [26.1, 29.8]; $Mdn = 20.0$; Mann–Whitney $U = 150582$, $p < .001$, $r_{tb} = .303$). Preference ties, while statistically different ($U = 34092$, $p = .022$, $r_{tb} = -.113$), exhibited small magnitude differences (Audio: $M = 3.9$; Text: $M = 3.3$ on a 100-point rating gap). We represent this as a cumulative distribution function in Figure 1.

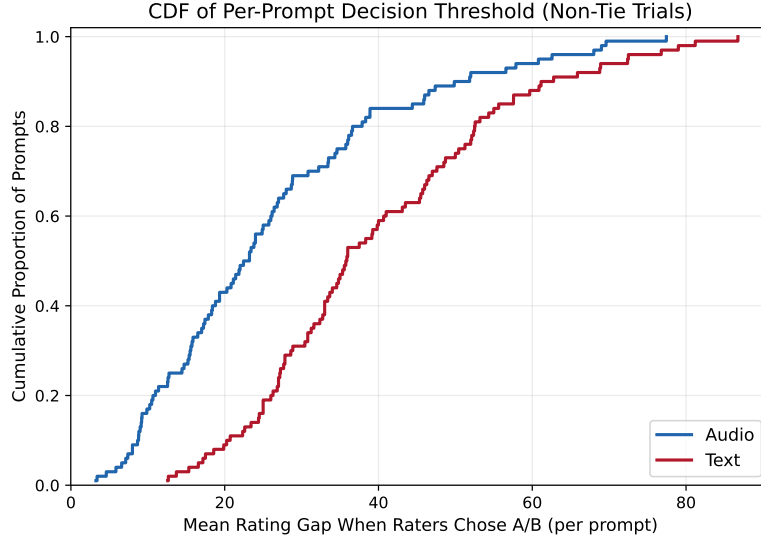


Figure 1: Cumulative distribution of the mean per-prompt rating gap on trials where raters declared a preference (A or B rather than Tie).

Table 1: Fixed effects from the linear mixed-effects model predicting raw ratings.

Predictor	b	95% CI	p
Intercept	74.80	[70.72, 78.87]	< .001
Presented first	-2.61	[-4.65, -0.58]	.012
Modality (text)	-6.21	[-11.28, -1.14]	.017
Character length (z)	3.35	[1.10, 5.60]	.004
Trial number (z)	0.03	[-1.02, 1.07]	.960
Presented first $\times$ Modality	-0.51	[-3.42, 2.40]	.732
Modality $\times$ Char. length	1.53	[0.03, 3.02]	.045
Modality $\times$ Trial number	0.90	[-0.59, 2.39]	.236

Note. Reference level for modality is audio. CIs are 95% bootstrap intervals.

Stimulus length, order, and fatigue. Modality moderated the effect of stimulus length ($b = 1.53$, 95% CI [0.03, 3.02], $p = .045$): longer responses predicted higher ratings in both modalities, but the effect was stronger for text (text: $b = 4.92$, $p < .001$; audio: $b = 3.52$, $p = .004$). A significant recency bias emerged, with second-presented items rated more favorably ($b = -2.61$, $p = .012$), without modality differences ($b = -0.51$, $p = .732$). Ratings remained stable across trials ($b = 0.03$, $p = .960$) with no modality-specific drift ($b = 0.90$, $p = .236$).

Practice Effects. A separate model on log-transformed trial duration showed completion times decreased over trials in both modalities (audio: $b = -0.09$, $p < .001$, $\sim 8.6\%$ reduction per SD of trial number; text: $b = -0.17$, $p < .001$, $\sim 16.0\%$), with text participants speeding up more ($b = -0.10$, $p < .001$). Despite these speedups, ratings remained stable, suggesting increased efficiency rather than declining effort.

4.2 Reliability of Audio Preference Data

Session Characteristics. Appendix Table 5 summarizes session characteristics. Audio trials took nearly twice as long as text trials ($M = 136$ vs. 79 sec; $U = 255.0$, $p < .001$). Attention check failure rates did not differ significantly between conditions (Audio: 13.2%; Text: 9.4%; Fisher’s exact test, $p = .761$), and participants rated both tasks as similarly difficult ($U = 992.0$, $p = .868$).

Inter-Rater Reliability (ICC). We estimated ICC(2,1) and ICC(2, k) via linear mixed-effects models with crossed random intercepts for stimuli and participants, with order/trial position as covariates

Table 2: Inter-rater reliability metrics by modality.

Metric	Audio	Text	p
Avg. raters per stimulus	9.2	8.9	—
ICC(2,1)	.333 [.286, .365]	.295 [.250, .325]	.127
ICC(2, k)	.821 [.787, .841]	.788 [.748, .811]	.113
Krippendorff’s α (cont.)	.315 [.230, .388]	.228 [.157, .298]	.069
Krippendorff’s α (ord.)	.157 [.109, .206]	.217 [.167, .268]	.119
Krippendorff’s α (nom.)	.010 [−.007, .026]	.031 [.012, .049]	.437

Note. Values are point estimates with 95% bootstrap CIs from stimulus resampling (2,000 for ICC; 5,000 for α). p -values from permutation tests of the Audio–Text difference.

and 95% CIs from 2,000 bootstrap resamples. Modality differences were tested by permutation (2,000 iterations, pseudo-Audio/Text groups preserving sample sizes). No significant differences were observed between ICC values for audio vs text (ICC(2,1) = .333 vs. .295, $p = .127$; ICC(2, k) = .821 vs. .788, $p = .113$).

Inter-Rater Reliability (Krippendorff’s α). We computed Krippendorff’s α for continuous ratings, ordinal preferences (A/Tie/B relative to presentation order, adjacent disagreements weighted less), and nominal preferences (raw A/Tie/B, all disagreements equal), with 95% CIs from 5,000 bootstrap resamples and modality tested by permutation (5,000 iterations). Unlike ICC, α was computed without covariate adjustment. No differences were significant (Table 2). Nominal α was near zero in both conditions, given the counterbalanced presentation order of our experimental design (i.e. agreement would assign opposite A/B labels).

4.3 Modality Effects on Preferences

Cross-Modality Agreement. For each prompt, we asked whether audio and text raters agreed on the preferred response. Within each modality, we averaged raw ratings across the ~ 9 raters per stimulus (justified by the good aggregate reliability, ICC(2, k) $\approx .80$) and identified the response with the higher mean. We then compared these per-prompt winners across modalities. At a zero-difference threshold, audio and text agreed on the winner for 53% of prompts (53/100; binomial vs. 50%, $p = .31$). Agreement rose with stricter decisiveness thresholds while the number of decisive pairs dropped (Appendix Figure 5); including ties-vs.-decisive as disagreements drops agreement until a threshold of ~ 10 .

Prompt-Specific Modality Effects. To test whether modality systematically shifted preferences, we computed per-trial preference scores (Clip₀ – Clip₁ ratings) and fit linear mixed-effects models with modality as a fixed effect and crossed random effects for prompt and participant. The fixed effect of modality was not significant ($\beta = 1.71$, $p = .458$), indicating no overall preference shift between modalities. However, adding random slopes for modality by prompt significantly improved model fit ($\chi^2(2) = 22.1$, $p < .001$; random slope $SD = 12.93$), indicating that modality affected preferences differently across prompts—some prompts elicited stronger preferences in audio, others in text.

4.4 AI-Human Alignment

Average AI-Human Rating Error. For each clip we computed mean AI and human ratings within modality and the absolute error $|AI - Human|$, then fit an LMM with modality as a fixed effect and clip as a random intercept.

PRISM vs. External Annotators. To test whether the original PRISM ratings, where the person interacting with the model did the rating, differed from external annotators, we conducted paired t -tests comparing PRISM ratings to each external group (AI-Text, AI-Audio, Human-Text, Human-Audio) at the clip level (raw score mean). We report our results in Table 3. Notably, human-audio ratings aligned most closely with PRISM ratings ($\bar{d} = 1.15$, $p = .374$), suggesting that audio evaluation may better approximate the original interactive rating context, though the non-significant result precludes strong claims.

Table 3: Model estimates for rating comparisons.

Comparison	$\bar{d}$	p
<i>AI-Human Absolute Error</i>		
Audio MAE	12.27	—
Text MAE	14.29	—
Text vs. Audio (LMM)	2.02	.035
<i>PRISM vs. External Annotators</i>		
vs. AI-Text	−1.84	.015
vs. AI-Audio	−9.18	< .001
vs. Human-Text	8.03	< .001
vs. Human-Audio	1.15	.374

Note. For AI-Human error, $\bar{d}$ is the mean absolute error (MAE) between AI and human clip-level ratings on a 1–100 scale; the LMM row tests the modality difference. For PRISM comparisons, $\bar{d}$ is the mean signed difference (PRISM – comparison group).

AI Difference Predicts Human Agreement. For each prompt we computed human ICC(2,1) and the absolute mean difference between AI ratings for the two responses, then regressed $\text{ICC} \sim \text{AI_diff} \times \text{Modality}$ with Spearman correlations as non-parametric checks (Table 4). When AI ratings sharply distinguished between responses, humans agreed more with each other in both modalities. Results were robust to removal of influential observations (Cook’s D).

Table 4: Regression: AI Differentiation and Human Agreement

Predictor	B	p
AI Differentiation	0.014	<.001
Modality (Text)	−0.010	.76
AI Diff $\times$ Modality	−0.003	.15
<i>Spearman correlations</i>		
Audio	$r = .23, p = .02$	
Text	$r = .38, p < .001$	

DV = per-prompt and modality ICC; $N = 100$

5 Discussion

Using a controlled cross-modal experiment, we provide critical analyses examining the reliability and consistency of human preference ratings across text and audio. While modalities are comparably reliable in aggregate, they produce meaningfully different evaluation dynamics: audio raters operated with narrower decision thresholds, showed reduced sensitivity to response length, and described their preferences in user-oriented rather than content-depth terms.

5.1 Audio Preference Reliability

Moderate reliability requires at least 3 raters per stimulus. To our knowledge, our work is the first to systematically characterize the number of raters needed to reach specific reliability thresholds in the preference annotation literature. As reviewed in §2.3, prior work has reported only single-rater pairwise agreement and not characterized how reliability scales with rater count. In line with Koo et al.’s [17] conservative ICC thresholds (Figure 2), agreement across single raters ($k = 1$) is *poor* $\text{ICC}(2,1) = .333$ audio, .295 text. At least three raters are required to achieve *moderate* consistency, while *good* agreement is observed with ~ 9 raters $\text{ICC}(2,k) = .821$ audio, .788 text, and diminishing returns beyond this. Practically, *moderate* reliability preference data demands a $3\times$ *increase* in annotation cost, while reaching *good* agreement demands a $9\times$ *increase* – which may or may not be justified depending on the requirements of the downstream task.

Assuming fixed (text) decision thresholds will misclassify audio preferences. The compressed audio margin (Figure 1) matters for pipelines that convert continuous ratings into binary labels: a fixed

binarization threshold calibrated on text may misclassify audio pairs as ties or reverse preferences, inflating label noise for reward model training. We recommend calibrating thresholds per modality or, where possible, retaining the continuous signal to preserve preference strength.

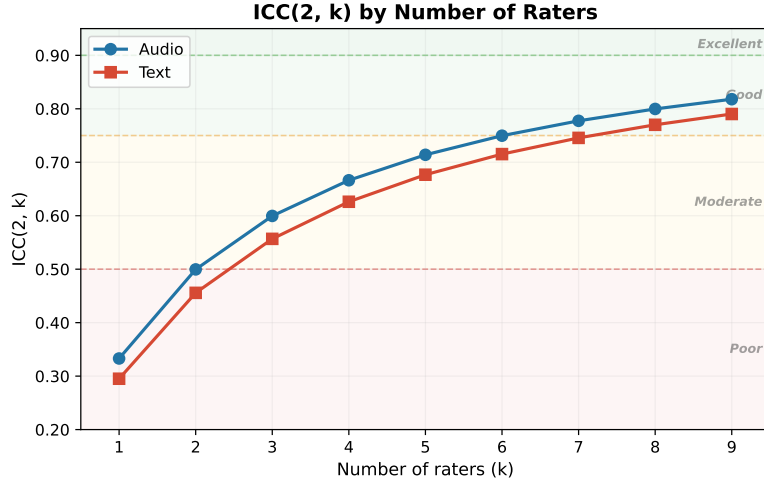


Figure 2: $ICC(2, k)$ as a function of the number of raters, derived from variance components of a crossed random-effects model. Shaded bands indicate conservative interpretation thresholds per [17]. The observed average number of raters per stimulus was $k \approx 9.2$ (audio) and $k \approx 8.9$ (text).

Recency Bias Must Be Handled Carefully. Text preferences are often presented side-by-side to mitigate order effects [2], but this is impossible for long audio. Our results confirm a strong recency bias that did not differ by modality, indicating it is a general problem of sequential presentation rather than audio-specific. Counterbalanced presentation order can mitigate this at the aggregate level but inflates individual-level disagreement, likely explaining our near-zero nominal Krippendorff’s α . Re-listening to the first clip after hearing both has also been proposed [30]; in our study, 214/920 audio trials (23.3%) involved voluntary re-listens to the first clip after the second. Binary ratings are especially susceptible to recency since order cannot be adjusted post-hoc, so we recommend continuous ratings for audio preference collection.

Longer Rating Sessions Are More Efficient Without Loss of Data Quality. Despite audio trials taking nearly twice as long, participants rated both conditions as equally difficult, suggesting the additional time reflects audio’s pacing rather than higher cognitive demand. Ratings remained stable across trials even as completion times decreased, indicating practice-driven efficiency rather than declining effort. While time-on-task effects are well documented in crowdsourcing [21, 19], no prior work has examined them for preference annotation specifically. Our results suggest sessions of up to one hour are viable for both modalities, though future work should identify the degradation boundary.

5.2 Cross-Modality Considerations

Text and Audio Converge Only When Preferences are Strong. Cross-modality agreement was only 53% at a zero threshold, rising with stricter decisiveness (Figure 5). The divergence is prompt-specific rather than a uniform directional shift: the fixed effect of modality was non-significant, but random slopes for modality-by-prompt substantially improved model fit, implying *modality can swing preferences by nearly half a standard deviation in either direction* depending on the prompt. Our TF-IDF analysis of written justifications offers one interpretation: audio raters more frequently referenced “user” and “help”, while text raters emphasized “detail” and “response” (Appendix G). This is consistent with audio evaluation promoting a holistic, user-oriented frame—*does this response serve the person?*—while text evaluation encourages analytic attention to content depth, aligning with evidence that users anthropomorphize audio interfaces more readily than text-based ones [9].

Audio Preferences are Less Susceptible to Length Bias. Length bias is a well-documented confound in text preference annotation [23], and we replicate it: longer text responses predicted a $\sim 40\%$ larger rating increase than longer audio responses (interaction $p = .045$). Audio’s fixed pacing likely

constrains length as a heuristic since listeners cannot skim or visually gauge response size. This extends Cho et al. [5], who showed verbosity is penalized more heavily in audio than text—though their best model still generated longer-than-baseline responses while being preferred, indicating audio length bias is moderated rather than reversed. One speculative use case: audio-derived labels could serve as a complementary signal to debias text-trained reward models by upweighting audio scores, at the cost of $\sim 2\times$ longer annotation per trial—whether this trade-off pays off remains open.

5.3 Synthetic Rating Applications

Synthetic ratings are an early identifier of high/low consistency prompts. When AI ratings distinguished between two responses, human raters also agreed more with each other, in both modalities. This finding has immediate practical value: *synthetic ratings can serve as a pre-screening signal to triage prompts before human annotation*. Prompts where AI models strongly differentiate likely need fewer human raters, while ambiguous AI ratings may warrant larger annotator pools to resolve genuine disagreement. Such adaptive sampling could substantially reduce annotation costs without sacrificing label quality. To our knowledge, this relationship between AI discriminability and human inter-rater reliability has not been previously reported, in part because ICC is rarely computed in the preference annotation literature. Future work may use our ICC data to validate other LLMs.

Modality-dependent AI Ratings Are Viable Human Proxies. AI-human absolute error was significantly higher for text than audio. While modest on a 100-point scale, this gap is notable given audio’s smaller decision thresholds—a 2-point discrepancy represents a proportionally larger share of the signal that distinguishes preferred from non-preferred audio responses. Prior work reports $\sim 80\%$ GPT-4o-human agreement for text [29], and RLAIIF can substitute AI for human labels in text with minimal performance loss [18]. Our results extend this to audio, suggesting synthetic ratings are viable proxies in both modalities, though tighter AI-human alignment in audio warrants further investigation into whether end-to-end audio models capture signals unavailable to text pipelines. However, these findings apply to a single (albeit popular) model family.

6 Limitations

Converting text-based interactions to audio via TTS is currently the standard approach for generating speech preference data. As such our findings speak to this common use-case, and allowed us to deliberately control for confounds by: utilizing a single static voice profile (af_heart), removing potentially harmful or offensive content, excluding “safety” and “refusal” scenarios, and a US-based, English-speaking sample of participants. However, TTS-rendered speech lacks paralinguistic nuances of natural speech—emotional prosody, hesitation, background noise—that may influence real-world preference judgments. At the same time speaker-characteristics (e.g. identity, gender, accent, pitch, culture) and non-innocuous/toxic semantic content are also likely to influence real-world preferences. Preference differences between non-Western populations and non-native English speakers are also largely unexplored. At the same time, additional context such as conversation history in text requires additional considerations for audio—having users listen to full conversations could be infeasible for high quality data collection, in the face of varying contextual depth. As such we interpret our findings, particularly around the number of annotators required to achieve *moderate* reliability as a lower bound relative to more naturalistic and unconstrained audio settings. These additional factors will likely require modified decision thresholds and additional annotators to account for increasing variance.

7 Conclusion

Despite the rapid development and deployment of speech audio models in recent years, differences between audio and text modalities has been understudied. We find that modality meaningfully shapes evaluation behavior: preferences vary by prompt rather than shifting uniformly, audio judgments show narrower decision thresholds and reduced length sensitivity, and synthetic ratings align with human judgments while predicting inter-rater reliability. Our work suggests that Text-derived TTS preferences cannot be *assumed* to generalize to speech, and modality should be treated as a first-class consideration in preference-based RL pipelines. At the same time, future work should systematically include more naturalistic audio factors and diverse speech conditions to further understand the differences between modalities.

Acknowledgments and Disclosure of Funding

Funding for this project was supported by the NIH National Library of Medicine’s T15 Biomedical Informatics and Data Science Research Training Program (Grant T15LM011271).

We’d like to thank Zachary Novack and Saketh Kasibatla for their feedback and continued support throughout this project.

References

- [1] Yuntao Bai, Andy Jones, Kamal Ndousse, Amanda Askell, Anna Chen, Nova DasSarma, Dawn Drain, Stanislav Fort, Deep Ganguli, Tom Henighan, et al. Training a helpful and harmless assistant with reinforcement learning from human feedback. *arXiv preprint arXiv:2204.05862*, 2022.
- [2] Shankha Basu and Krishna Savani. Choosing among options presented sequentially or simultaneously. *Current Directions in Psychological Science*, 28(1):97–101, 2019.
- [3] Aaron Broukhim, Yiran Shen, Prithviraj Ammanabrolu, and Nadir Weibel. Preference-based learning in audio applications: A systematic analysis. *arXiv preprint arXiv:2511.13936*, 2025.
- [4] Shreyas Chaudhari, Pranjal Aggarwal, Vishvak Murahari, Tanmay Rajpurohit, Ashwin Kalyan, Karthik Narasimhan, Ameet Deshpande, and Bruno Castro da Silva. Rlhf deciphered: A critical analysis of reinforcement learning from human feedback for llms. *ACM Computing Surveys*, 58(2):1–37, 2025.
- [5] Hyundong Justin Cho, Nicolaas Paul Jedema, Leonardo F. R. Ribeiro, Karishma Sharma, Pedro Szekely, Alessandro Moschitti, Ruben Janssen, and Jonathan May. Speechworthy instruction-tuned language models. In Yaser Al-Onaizan, Mohit Bansal, and Yun-Nung Chen, editors, *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, pages 10652–10670, Miami, Florida, USA, November 2024. Association for Computational Linguistics.
- [6] Paul F Christiano, Jan Leike, Tom Brown, Miljan Martic, Shane Legg, and Dario Amodei. Deep reinforcement learning from human preferences. *Advances in neural information processing systems*, 30, 2017.
- [7] Yunfei Chu, Jin Xu, Qian Yang, Haojie Wei, Xipin Wei, Zhifang Guo, Yichong Leng, Yuanjun Lv, Jinzheng He, Junyang Lin, et al. Qwen2-audio technical report. *arXiv preprint arXiv:2407.10759*, 2024.
- [8] Geoffrey Cideron, Sertan Girgin, Mauro Verzetti, Damien Vincent, Matej Kastelic, Zalan Borsos, Brian McWilliams, Victor Ungureanu, Olivier Bachem, Olivier Pietquin, Matthieu Geist, Léonard Hussenot, Neil Zeghidour, and Andrea Agostinelli. Musicrl: aligning music generation to human preferences. In *Proceedings of the 41st International Conference on Machine Learning, ICML’24*. JMLR.org, 2024.
- [9] Michelle Cohn, Mahima Pushkarna, Gbolahan O. Olanubi, Joseph M. Moran, Daniel Padgett, Zion Mengesha, and Courtney Heldreth. Believing anthropomorphism: Examining the role of anthropomorphic cues on trust in large language models. In *Extended Abstracts of the CHI Conference on Human Factors in Computing Systems, CHI EA ’24*, New York, NY, USA, 2024. Association for Computing Machinery.
- [10] Qingkai Fang, Shoutao Guo, Yan Zhou, Zhengrui Ma, Shaolei Zhang, and Yang Feng. Llama-omni: Seamless speech interaction with large language models. *arXiv preprint arXiv:2409.06666*, 2024.
- [11] Leo Gao, John Schulman, and Jacob Hilton. Scaling laws for reward model overoptimization. In *International Conference on Machine Learning*, pages 10835–10866. PMLR, 2023.
- [12] Xiaoxue Gao, Chen Zhang, Yiming Chen, Huayun Zhang, and Nancy F Chen. Emo-dpo: Controllable emotional speech synthesis through direct preference optimization. In *ICASSP 2025-2025 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 1–5. IEEE, 2025.
- [13] hexgrad. Kokoro-82m: An open-weight text-to-speech model. <https://huggingface.co/hexgrad/Kokoro-82M>, 2025. Apache 2.0 License.
- [14] Ailin Huang, Boyong Wu, Bruce Wang, Chao Yan, Chen Hu, Chengli Feng, Fei Tian, Feiyu Shen, Jingbei Li, Mingrui Chen, et al. Step-audio: Unified understanding and generation in intelligent speech interaction. *arXiv preprint arXiv:2502.11946*, 2025.
- [15] Raisa Islam and Owana Marzia Moushi. Gpt-4o: The cutting-edge advancement in multimodal llm. *Authorea Preprints*, 2024.

- [16] Hannah Rose Kirk, Alexander Whitefield, Paul Rottger, Andrew M Bean, Katerina Margatina, Rafael Mosquera-Gomez, Juan Ciro, Max Bartolo, Adina Williams, He He, et al. The prism alignment dataset: What participatory, representative and individualised human feedback reveals about the subjective and multicultural alignment of large language models. *Advances in Neural Information Processing Systems*, 37:105236–105344, 2024.
- [17] Terry K Koo and Mae Y Li. A guideline of selecting and reporting intraclass correlation coefficients for reliability research. *Journal of chiropractic medicine*, 15(2):155–163, 2016.
- [18] Harrison Lee, Samrat Phatale, Hassan Mansoor, Thomas Mesnard, Johan Ferret, Kellie Lu, Colton Bishop, Ethan Hall, Victor Carbune, Abhinav Rastogi, and Sushant Prakash. Rlaif vs. rlhf: scaling reinforcement learning from human feedback with ai feedback. In *Proceedings of the 41st International Conference on Machine Learning*, ICML’24. JMLR.org, 2024.
- [19] Winter Mason and Siddharth Suri. Conducting behavioral research on amazon’s mechanical turk. *Behavior research methods*, 44(1):1–23, 2012.
- [20] Prolific. Prolific (version: May 2025). <https://www.prolific.com>, 2024. London, UK. First released in 2014. Accessed May 2025.
- [21] Judi E See, Steven R Howe, Joel S Warm, and William N Dember. Meta-analysis of the sensitivity decrement in vigilance. *Psychological bulletin*, 117(2):230, 1995.
- [22] Wei Shen, Rui Zheng, Wenyu Zhan, Jun Zhao, Shihan Dou, Tao Gui, Qi Zhang, and Xuan-Jing Huang. Loose lips sink ships: Mitigating length bias in reinforcement learning from human feedback. In *Findings of the Association for Computational Linguistics: EMNLP 2023*, pages 2859–2873, 2023.
- [23] Prasann Singhal, Tanya Goyal, Jiacheng Xu, and Greg Durrett. A long way to go: Investigating length correlations in rlhf. *arXiv preprint arXiv:2310.03716*, 2023.
- [24] Nisan Stiennon, Long Ouyang, Jeffrey Wu, Daniel Ziegler, Ryan Lowe, Chelsea Voss, Alec Radford, Dario Amodei, and Paul F Christiano. Learning to summarize with human feedback. *Advances in neural information processing systems*, 33:3008–3021, 2020.
- [25] TTS-AGI. Tts arena v2. <https://huggingface.co/spaces/TTS-AGI/TTS-Arena-V2>, 2024.
- [26] Peiyi Wang, Lei Li, Liang Chen, Zefan Cai, Dawei Zhu, Binghui Lin, Yunbo Cao, Lingpeng Kong, Qi Liu, Tianyu Liu, and Zhifang Sui. Large language models are not fair evaluators. In Lun-Wei Ku, Andre Martins, and Vivek Srikumar, editors, *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 9440–9450, Bangkok, Thailand, August 2024. Association for Computational Linguistics.
- [27] Zeqiu Wu, Yushi Hu, Weijia Shi, Nouha Dziri, Alane Suhr, Prithviraj Ammanabrolu, Noah A Smith, Mari Ostendorf, and Hannaneh Hajishirzi. Fine-grained human feedback gives better rewards for language model training. *Advances in Neural Information Processing Systems*, 36:59008–59033, 2023.
- [28] Dong Zhang, Zhaowei Li, Shimin Li, Xin Zhang, Pengyu Wang, Yaqian Zhou, and Xipeng Qiu. Speechalign: Aligning speech generation to human preferences. *Advances in Neural Information Processing Systems*, 37:50343–50360, 2024.
- [29] Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanghao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric Xing, et al. Judging llm-as-a-judge with mt-bench and chatbot arena. *Advances in neural information processing systems*, 36:46595–46623, 2023.
- [30] Slawomir Zielinski, Francis Rumsey, and Søren Bech. On some biases encountered in modern audio quality listening tests-a review. *Journal of the Audio Engineering Society*, 56(6):427–451, 2008.

A Broader Impacts

Our methodological findings and released audio corpus aim to support audio-aligned speech systems, an area where alignment research remains underdeveloped relative to text. Three considerations warrant care. First, better audio preference pipelines could in principle be misused to optimize persuasive or manipulative speech systems, though our contribution is methodological rather than a deployable model, making this path indirect. Second, our US-based, predominantly White, predominantly native-English-speaking sample (Appendix E) means reward signals derived from data collected via our protocol could entrench preferences that systematically devalue non-Western speech

patterns, accents, or dialects; we recommend demographically broader rater pools for any downstream deployment. Third, we excluded potentially harmful or offensive content from our stimuli, so our reliability and modality-effect results do not characterize safety or refusal scenarios and should not be generalized to safety-critical alignment contexts. We flag the latter two caveats in Limitations (Section 6).

B Study Interfaces

Figures 3 and 4 show the rater-facing interfaces for the audio and text conditions, respectively. The two interfaces were matched on layout, conversation history presentation, and rating instruments; they differ only in (i) whether responses were delivered as audio clips with playback controls or as rendered text, and (ii) the presence of an additional audio-quality slider in the audio condition. Both conditions used continuous 1–100 visual analog scales with hidden numeric values, followed by a discrete preference judgment (A / Tie / B) and an optional free-text justification.

Trial 3 of 20

User: Hi, can you explain quantum mechanics to me like I am a 5-year old?

Audio A

Play Replay -10s

0:00 0:17

Response Rating:

Terrible Perfect

Audio Quality Rating:

Terrible Perfect

Audio B

Play Replay -10s

0:00 0:43

Response Rating:

Terrible Perfect

Audio Quality Rating:

Terrible Perfect

Overall, which response do you prefer?

A Tie B

Why did you prefer one response over the other? (Optional)

Please share your reasoning...

Submit Ratings Exit Study Early

Please listen to both audio clips in full and adjust the rating slider for all ratings before submitting.

Figure 3: Interface for the Audio Task of preference selection. It includes conversation history, most recent query and two responses. Raters must rate each response from 1-100 on response quality (the content of the answer) and audio quality (the clarity and naturalness of speech). We also collect discrete preferences with the option to tie and optional free-response reasoning justification.

C Session and Trial Characteristics

Table 5 summarizes session- and trial-level descriptive statistics by modality. Audio sessions took substantially longer than text sessions both per trial (136 vs. 79 seconds) and overall (47 vs. 27 minutes), reflecting audio’s fixed temporal pacing rather than greater task difficulty: self-reported difficulty ratings were indistinguishable across conditions ($M = 2.50$ vs. 2.45 on a 5-point scale, $p = .868$). Completion and attention-check failure rates were also comparable between conditions, suggesting the longer audio sessions did not disproportionately tax participants or induce dropout. Together, these characteristics support the interpretation in §5.1 that audio annotation imposes a time cost without a corresponding effort cost.

Trial 3 of 20

User: Hi, can you explain quantum mechanics to me like I am a 5-year old?

Response A

Imagine the world is made of tiny LEGO blocks called atoms. Some blocks like to stick together, and some don't. Quantum mechanics is the special rulebook for the very, very tiny LEGO blocks. The fun part is that these tiny blocks can sometimes be in two places at once, like if your favorite toy could be on the table and on the floor at the same time until you look at it!

Response Rating:

Terrible

Perfect

Response B

Quantum mechanics is the branch of physics that describes how particles behave at very small scales. At those scales, things don't always act like the everyday objects you see — a particle can act like a wave, and we can't always know exactly where it is and how fast it's moving at the same time. Scientists use math to predict the chances of where a particle might be, instead of saying for sure.

Response Rating:

Terrible

Perfect

Overall, which response do you prefer?

A

Tie

B

Why did you prefer one response over the other? (Optional)

Please share your reasoning...

Submit Ratings

Exit Study Early

Please adjust the rating slider for all ratings before submitting.

Figure 4: Interface for the Text Task of preference selection. It includes conversation history, most recent query and two responses. Raters must rate each response from 1-100 on response quality (the content of the answer). We also collect discrete preferences with the option to tie and optional free-response reasoning justification.

Table 5: Session and trial characteristics by condition.

	Text	Audio	<i>p</i>
Completed (%)	92.5	96.2	.678 ^a
Session duration, min	27.0 (13.3)	47.3 (12.4)	< .001 ^b
Trial duration, sec	78.8 (41.6)	136.4 (35.9)	< .001 ^b
Difficulty rating	2.45 (0.90)	2.50 (0.96)	.868 ^b
Failed attention (%)	9.4	13.2	.761 ^a

Note. Values are *M* (*SD*) unless noted. ^aFisher's exact test. ^bMann-Whitney *U*.

D Cross-Modality Agreement by Decisiveness Threshold

Figure 5 extends the Section 4.3 analysis by plotting cross-modality winner agreement as a function of the minimum rating gap required to count a prompt as “decisive.” The blue curve restricts the comparison to prompts where both modalities produced a decisive winner at the given threshold; agreement rises monotonically from 53% at a zero-difference threshold to substantially higher rates once only strongly preferred pairs are retained, but the number of qualifying prompts (green bars) drops accordingly. The red curve treats tie-versus-decisive mismatches as disagreements, and consequently declines until the threshold reaches roughly 10 points—the point at which most prompts in both modalities have settled into clear preferences. The two curves together indicate that text and audio raters converge primarily on the prompts where preferences are unambiguous, and diverge on the ambiguous middle, consistent with the prompt-specific modality effects reported in Section 4.3.

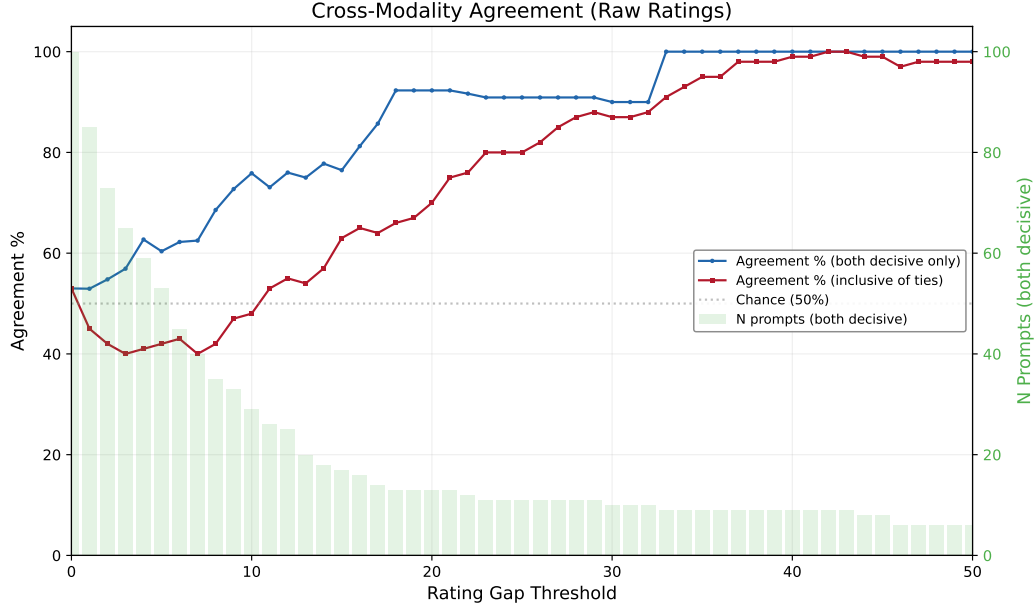


Figure 5: Cross-modality agreement between audio and text evaluations as a function of raw rating gap threshold. The blue line shows agreement percentage when both modalities produce a decisive winner (excluding ties), while the red line shows agreement when tie-versus-decisive outcomes are counted as disagreements. Green bars indicate the number of prompts where both modalities had a decisive winner at each threshold.

E Participant Demographics

Table 6 reports the demographic composition of the final analytic sample ($N = 94$) by modality condition. We report these descriptives both for transparency and to support the robustness check described below.

Participant demographics were largely comparable across the two modality conditions. Both groups were predominantly White (Text: 60.4%, Audio: 65.2%), non-Hispanic (91.7%, 84.8%), and native English speakers (85.4%, 91.3%). Text-condition participants were somewhat older on average ($M = 45.00$, $SD = 13.77$) than audio-condition participants ($M = 40.17$, $SD = 11.70$), and held bachelor’s degrees or higher at a higher rate (60.4% vs. 43.5%). As a robustness check, we re-fit the omnibus model with education (bachelor’s+ vs. below) as an additional fixed effect; the modality effects reported in the main text were unchanged. Because participants were randomly assigned to modality, these differences reflect sampling variability rather than systematic selection.

F Robustness to Perceived Audio Quality

Because the audio condition introduced a fidelity dimension absent in text, we re-fit the audio-only simple-effects model with z -scored audio-quality ratings added as a covariate ($n = 1,840$ observations, 46 participants, 200 stimuli). Perceived audio quality was, unsurprisingly, a strong positive predictor of response ratings ($b = 7.77$, 95% CI [6.51, 9.03], $p < 10^{-31}$): a 1-SD increase in audio-quality rating corresponded to a ~ 7.8 -point increase in response rating. However, the substantive effects from the main text were preserved. The recency bias remained significant ($b = -2.28$, 95% CI [-4.11, -0.44], $p = .015$; original $b = -2.63$, $p = .006$, $\sim 14\%$ attenuation). The length effect was essentially unchanged ($b = 3.27$, 95% CI [1.05, 5.50], $p = .004$; original $b = 3.53$, $p = .004$). The trial-number effect remained null ($b = 0.28$, $p = .57$; original $b = 0.02$, $p = .96$). Audio quality therefore contributes meaningfully to response-quality judgments but does not account for the order or length effects on which our cross-modal results rest.

Table 6: Participant Demographics by Condition

	Text ($n = 48$)	Audio ($n = 46$)	Total ($N = 94$)
<i>Age (years)</i>			
<i>M (SD)</i>	45.00 (13.77)	40.17 (11.70)	42.64 (12.96)
Median	44	40	41
Range	22–75	20–72	20–75
<i>Gender</i>			
Woman	31 (64.6%)	25 (54.3%)	56 (59.6%)
Man	17 (35.4%)	20 (43.5%)	37 (39.4%)
Agender	0 (0.0%)	1 (2.2%)	1 (1.1%)
<i>Race</i>			
White	29 (60.4%)	30 (65.2%)	59 (62.8%)
Black or African American	7 (14.6%)	10 (21.7%)	17 (18.1%)
Asian	7 (14.6%)	2 (4.3%)	9 (9.6%)
Two or More Races	3 (6.2%)	3 (6.5%)	6 (6.4%)
American Indian/Alaska Native	2 (4.2%)	0 (0.0%)	2 (2.1%)
Prefer not to say	0 (0.0%)	1 (2.2%)	1 (1.1%)
<i>Ethnicity</i>			
Not Hispanic or Latino	44 (91.7%)	39 (84.8%)	83 (88.3%)
Hispanic or Latino	3 (6.2%)	7 (15.2%)	10 (10.6%)
Prefer not to say	1 (2.1%)	0 (0.0%)	1 (1.1%)
<i>English Proficiency</i>			
Native speaker	41 (85.4%)	42 (91.3%)	83 (88.3%)
Fluent	6 (12.5%)	4 (8.7%)	10 (10.6%)
Conversational	1 (2.1%)	0 (0.0%)	1 (1.1%)
<i>Education</i>			
Bachelor’s degree	21 (43.8%)	14 (30.4%)	35 (37.2%)
Some college, no degree	7 (14.6%)	13 (28.3%)	20 (21.3%)
High school or equivalent	8 (16.7%)	8 (17.4%)	16 (17.0%)
Master’s degree	6 (12.5%)	5 (10.9%)	11 (11.7%)
Associate degree	4 (8.3%)	5 (10.9%)	9 (9.6%)
Doctorate	2 (4.2%)	1 (2.2%)	3 (3.2%)
Bachelor’s+ (%)	60.4	43.5	52.1

Note. All participants resided in the United States. Demographics reflect participants who passed the attention check.

G Qualitative Analysis

To compare how participants explained their preferences across modalities, we analyzed 658 written justifications (347 audio, 311 text) using TF-IDF with Porter stemming (vocabulary = 223 terms). For each term, we tested whether document frequency differed by modality using χ^2 tests on presence/absence, with Benjamini-Hochberg FDR correction ($q = .05$). Excluding modality-specific terms, audio justifications more frequently mentioned “user” ($\chi^2 = 12.8$, $p < .001$) and “help” ($\chi^2 = 9.8$, $p = .002$), while text justifications more frequently mentioned “response” ($\chi^2 = 14.5$, $p < .001$), “detail” ($\chi^2 = 14.5$, $p < .001$), and “give” ($\chi^2 = 17.0$, $p < .001$), suggesting audio raters focused more on whether the response served the user’s needs while text raters attended more to the depth and specificity of the response content.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper’s contributions and scope?

Answer: [\[Yes\]](#)

Justification: This work is a step towards best practices in preferences data collection for both text and audio modalities. We focus on practical guidelines and novel understandings of how current speech generation pipelines may be different than text pipelines and how to collect higher quality data cross-modality.

Guidelines:

- The answer [\[N/A\]](#) means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A [\[No\]](#) or [\[N/A\]](#) answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: We explicitly state that our findings related to speech are limited in that they are conducted on a single TTS voice and further works should investigate this further. Additionally, our synthetic data reflects GPT-4o and GPT-4o-Audio-Preview and additional works are needed to generalize these claims.

Guidelines:

- The answer [\[N/A\]](#) means that the paper has no limitation while the answer [\[No\]](#) means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate “Limitations” section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren’t acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [N/A]

Justification: This paper's focus is on practical considerations of data collection for speech and text preferences.

Guidelines:

- The answer [N/A] means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: We fully outline our methods, in detail, such that others can convert the PRISM data to speech using kokoro and running a similar study using our code or a similar interface to the image presented in the paper.

Guidelines:

- The answer [N/A] means that the paper does not include experiments.
- If the paper includes experiments, a [No] answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).

- (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: We make both our data and code available to reproduce our study plus provide instructions on how to set up similar studies.

Guidelines:

- The answer [N/A] means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://neurips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so [No] is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://neurips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer) necessary to understand the results?

Answer: [N/A]

Justification: Models were not trained in this study, but analysis and data collection methods are clearly articulated.

Guidelines:

- The answer [N/A] means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: We report appropriate p-values and confidence intervals where relevant, and make statistical adjustments where relevant (Benjamini-Hochberg FDR correction in qualitative analysis)

Guidelines:

- The answer [N/A] means that the paper does not include experiments.
- The authors should answer [Yes] if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g., negative error rates).
- If error bars are reported in tables or plots, the authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: All analyses were run on a standard laptop CPU; no GPU or cluster compute was required. Bootstrap analyses (2,000–5,000 iterations) completed in minutes.

Guidelines:

- The answer [N/A] means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: Data is collected and anonymized. Data was identified as minimal risk by our institution IRB.

Guidelines:

- The answer [N/A] means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer [No], they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: Discussed in Appendix A; see also Limitations (Section 6) and Appendix E.

Guidelines:

- The answer [N/A] means that there is no societal impact of the work performed.
- If the authors answer [N/A] or [No], they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate Deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pre-trained language models, image generators, or scraped datasets)?

Answer: [Yes]

Justification: Our released assets consist of TTS-converted audio of the publicly available PRISM dataset’s unguided conversations, which were screened to remove offensive, discriminatory, or emotionally distressing content. We do not release pretrained models, and the data poses no meaningful misuse or dual-use risk beyond what already exists in the source PRISM corpus.

Guidelines:

- The answer [N/A] means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: We cite all assets used: the PRISM dataset (Kirk et al., 2024), the Kokoro TTS model (hexgrad), TTS-Arena benchmark, GPT-4o and GPT-4o-Audio-Preview (OpenAI

APIs), and Prolific for participant recruitment. We use these assets in accordance with their respective licenses and terms of service.

Guidelines:

- The answer [N/A] means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset’s creators.

13. New assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [\[Yes\]](#)

Justification: We release the 3,113 TTS-converted audio conversations, data annotations, interface code, and our analysis code with documentation. The dataset is derived from the publicly released PRISM corpus, whose participants consented to public release of their conversations. Documentation of selection, removals, and replacements is provided alongside the released data.

Guidelines:

- The answer [N/A] means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and research with human subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [\[Yes\]](#)

Justification: We recruited 106 participants from Prolific and compensated them at \$12/hour, above the US federal minimum wage. The study interfaces, including all instructions and rating instruments shown to participants, are presented in Appendix Figures 3 and 4. The demographic questionnaire, attention check procedure, exclusion criteria, and task structure are described in Section 3.2, with full participant demographics reported in Appendix Table 6.

Guidelines:

- The answer [N/A] means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.

- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional review board (IRB) approvals or equivalent for research with human subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [\[Yes\]](#)

Justification: The study was deemed minimal risk by our institutions review board. We report this explicitly in the paper.

Guidelines:

- The answer [\[N/A\]](#) means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. Declaration of LLM usage

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does *not* impact the core methodology, scientific rigor, or originality of the research, declaration is not required.

Answer: [\[Yes\]](#)

Justification: While LLMs were used for manuscript writing support, they did not play a significant role in methods development. GPT-4o and GPT-4o-Audio-Preview were used to generate synthetic ratings analyzed against human ratings, with default API parameters

Guidelines:

- The answer [\[N/A\]](#) means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy in the NeurIPS handbook for what should or should not be described.