

**NATIVE VERSUS NON-NATIVE SPEECH PERCEPTION IN NOISE: EFFECT OF
REVERBERATION AND SIGNAL TO NOISE RATIO ON LISTENING EFFORT**

SHIVANI SHARMA

Registration No: P01II24S023030

A Dissertation Submitted in Part of Fulfilment of the

Degree of Master of Science

(Audiology)

University of Mysore



ALL INDIA INSTITUTE OF SPEECH AND HEARING,

MANASAGANGOTHRI, MYSURU -570006

JUNE, 2026

CERTIFICATE

This is to certify that this dissertation entitled " **Native versus Non-Native Speech Perception in Noise: Effect of Reverberation and Signal to Noise Ratio on Listening Effort**" is a bonafide work submitted as part of fulfilment of the degree of Master of Science (**Audiology**) of the student with Registration Number **P011124S023030**. This has been carried out under the guidance of a faculty of this institute and has not been submitted earlier to any other university for the award of any other Diploma or Degree.

Mysuru
8th June, 2026



Dr. M. Pushpavathi

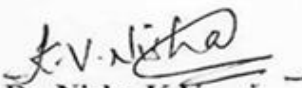
Director

All India Institute of Speech and Hearing
Manasagangothri, Mysuru-570 006

CERTIFICATE

This is to certify that this dissertation entitled " **Native versus Non-Native Speech Perception in Noise: Effect of Reverberation and Signal to Noise Ratio on Listening Effort** " has been prepared under my supervision and guidance. It is also being certified that this dissertation has not been submitted earlier to any other university for the award of any other Diploma or Degree.

Mysuru
June, 2026


Dr. Nisha K V

Guide

Scientist B,

COE – Centre for Hearing Sciences,
Department of Audiology,
All India Institute of Speech and Hearing,
Manasagangothri Mysuru – 570 006



Dr. P. Arivudai Nambi

Co-Guide

Associate Professor of Audiology,
COE – Centre for Hearing Sciences,
Department of Audiology,
All India Institute of Speech and Hearing,
Manasagangothri Mysuru – 570 006

DECLARATION

This is to certify that this dissertation entitled "**Native versus Non-Native Speech Perception in Noise: Effect of Reverberation and Signal to Noise Ratio on Listening Effort**" is the result of my own study under the guidance of Dr. Nisha K V, Scientist B, COE-Centre for Hearing Sciences, Department of Audiology, All India Institute of Speech and Hearing, Manasagangothri, Mysuru and Dr. P. Arivudai Nambi, Associate Professor of Audiology, COE – Centre for Hearing Sciences, Department of Audiology, All India Institute of Speech and Hearing, Mysuru and has not been submitted earlier to any other university for the award of any other Diploma or Degree.



Mysuru
June, 2026

Registration no. P01II24S023030

Acknowledgements

*To my **Bholenath** — everything begins and ends with you, whose divine blessings, strength, and guidance have carried me through every step of this journey. On the days I felt lost, you were the way. Thank You, my Shiva, for being my strength, my calm, and my constant, for being the foundation beneath everything I am and everything I do.*

There are moments in life that quietly shape everything — and completing this dissertation has been one of them. This work carries within it the fingerprints of so many souls, and I am overwhelmed with gratitude trying to find the right words.

*I am deeply grateful to **Dr. M. Pushpavathi**, Director of AIISH, and **Dr. Niraj Kumar.**, HOD Audiology, for providing the platform and institutional support that made this research possible.*

*First and foremost, I bow in deep gratitude to **Dr. Nisha K.V.**, my guide and mentor, whose wisdom, patience, and unwavering belief in me made this journey possible. Your dedication to your work is something I have witnessed firsthand and with an excellence that quietly sets the standard for everyone around you. You saw potential in me even when I could not see it in myself. Thank you for pushing me to think harder, write clearer, and never settle for less than my best and I carry your guidance with me as one of the greatest gifts this academic journey has given me. You did not just supervise a dissertation — you helped build a researcher. Thank you, Ma'am, from the very bottom of my heart.*

*To **Dr. P. Arivudai Nambi**, my co-guide. Your brilliance is the kind that makes difficult things look effortless, and working with you has been one of the greatest privileges of this journey. You brought the stimulus and the statistical models to life with so much precision and care. Thank you, Sir, from the bottom of my heart.*

To the ninety participants (seniors, juniors and my batchmates) who gave their time, their patience, and their voices to this study — this work exists because of you. You sat through hours of listening tasks with grace and goodwill, and I will never forget that. This dissertation is, at its heart, about the effort of listening. You lived that effort, and I am deeply grateful.

*To my beloved family, you are the ground beneath my feet. To my **Papa and Mummy** — my foundation, my greatest blessing, who sacrificed more than I will ever fully know, and who never once asked me to stop dreaming You motivated me when I was lost, encouraged me when I was scared, and pushed me gently forward every single time I wanted to stop. I would not have crossed this finish line without you. This dissertation has my name on it, but it was built on your love, your sacrifices, and your unshakeable belief in me. I hope I have made you proud, because everything I am, I owe to you both. This one is yours.*

To my siblings, **Pummi di and Om** (my babu, my baby brother), my absolute treasures who gets me, gets exactly what I need without me having to say a word, who tolerated my frustration and stress with extraordinary patience, and who always knew when I needed laughter more than advice and have this ridiculous gift of making me laugh, the kind that makes everything feel lighter — at exactly the right moment — thank you.

A very special note of gratitude goes to my six-month-old niece, my **Gullu** whose presence brought immense joy and comfort during one of the most demanding phases of this dissertation. In moments of stress, exhaustion, and self-doubt, her innocent smile, cheerful laughter, and bright little expressions had a remarkable way of lifting my spirits and bringing a sense of peace and happiness. Though she may never realize it, her presence illuminated my days, reminded me to pause and appreciate life's simple joys, and provided a source of emotional strength throughout this journey. She truly became a little ray of sunshine who made this challenging period much brighter and more meaningful.

To **Roushan Sir** — my senior, my constant support, my pillar of strength and so much more than I could ever put into words on this page. You were there through more of this journey than anyone realises, in ways that only we know. Thank you for the encouragement, the patience, the calm you brought into my chaos, the conversations that somehow made everything feel possible again, and for believing in me.

To my wonderful friends — **Shambhavi, Shaily, Shubhangi, Anmol, Lucky, Apurva, Divya, Mamta, Bhumika, Abu, Sumanyu, Srijita, Vikas and Prasanna**. Having friends like you is not something everyone gets, and I do not take it for granted for even a second. You were the laughter in the middle of the stress, the distraction when I needed it most, and the warmth that made even the hardest days feel bearable. Thank you for the memories, the madness, and the love. I am so incredibly lucky to have every single one of you. This chapter of my life has your names written all over it and special thanks to **Shambhavi** — my bestie, my person, my mood-fixer. There is something truly special about having someone you can share everything with — the good days, the bad days, the overthinking, the breakthroughs, the absolute chaos — and you have been that person for me through all of it. Every single time we talked, without fail, you lifted me up. Thank you for listening, for laughing with me, for never getting tired of my dissertation stress, and for always, always making me feel better. I am so glad you are mine.

To my friends, who became family along the way (my Sakhi Saheli group-**Varsha, Swasti, Sonal, Shruti**). Thank you for the chai, the conversations, the parties and the pep talk at unreasonable hours. You people have made me feel like home in this journey and carried me more than you know.

To my seniors, who made this journey so much easier.

To **Ananya Maam** — I lost count of how many times I came running to you with a problem, a confusion, or a panic, and every single time you sat with me and explained everything so calmly and so clearly until it all made sense. You never once made me feel silly for asking. You have this beautiful way of making difficult things feel approachable, and your patience and warmth made such a difference to me. I am so grateful you were there through all of it.

To **Bharani Maam** — thank you for teaching me to use and run the Paradigm software with such care and thoroughness. Without your guidance, an entire pillar of this research would not have stood. You took the time to make sure I truly understood, and that meant everything.

*To **Suraj Sir** — thank you for helping me out whether it was installing SPSS at an odd hour, troubleshooting calibration, or just figuring out what to do next — you showed up every single time without making me feel like a. Thank you for being so dependably, wonderfully helpful.*

To all three of you — thank you for being the kind of seniors that juniors feel lucky to have.

To every person who crossed my path during this journey — those whose names fill these pages and those who never knew the role they played — thank you. As I turn the last page of this chapter, I find myself not just holding a completed dissertation, but carrying something far greater — a richer version of myself, shaped by every challenge that broke me a little and every person who helped put me back together. The completion of this dissertation would not have been possible without the guidance, encouragement, support, and kindness of so many individuals who stood by me throughout this process. This journey has been both challenging and rewarding, shaping me not only academically but also personally. I am deeply thankful for every lesson learned, every obstacle overcome, and every person who helped make this achievement possible. This journey gave me more than a degree. It gave me depth, perspective, humility, and a gratitude so profound it sits in my chest like something permanent. I do not step forward from here unchanged. I step forward braver, softer, stronger, and endlessly thankful — with my eyes wide open and my heart even wider. I carry forward this accomplishment with immense gratitude and humility.

Har Har Mahadev.

TABLE OF CONTENTS

Chapter	Title	Page Number
	Abstract	1
1	Introduction	3
2	Review of Literature	12
3	Methods	28
4	Results	55
5	Discussion	136
6	Summary and Conclusions	164
	References	169
	Appendix	185

LIST OF TABLES

Table No.	Title	Page Number
2.1	<i>Classification of listening effort assessment methods showing physiological, cognitive/behavioral, and subjective rating measures</i>	23
3.1	<i>Demographic, audiological, cognitive, and language characteristics of the participants</i>	30
3.2	<i>Equipment/software involved in the adaptations related to stimulus generation and presentation of dual task paradigm</i>	34
3.3	<i>Pilot testing performance scores for primary and secondary tasks across participants</i>	40
4.1	<i>Fixed-effect estimates from the Beta GLMM showing the effects of group, RT, SNR on primary task scores</i>	59
4.2	<i>ANOVA results (Type II Wald χ^2 tests) for primary task performance</i>	62
4.3	<i>Pairwise group comparisons on primary task performance across RT and SNR conditions</i>	67
4.4	<i>Pairwise comparison of RT conditions on primary task performance</i>	70
4.5	<i>Pairwise comparison of SNR conditions (0 dB vs 4 dB) on primary task performance</i>	74
4.6	<i>Fixed-effect estimates from the LMER model showing the effects of group, RT, SNR on secondary task scores</i>	78
4.7	<i>ANOVA results (Type II Wald χ^2 tests) for secondary task performance</i>	81
4.8	<i>Pairwise group comparisons on secondary task performance across RT and SNR conditions</i>	89
4.9	<i>Pairwise comparison of RT conditions on secondary task performance</i>	94
4.10	<i>Pairwise comparison of SNR conditions (0 dB vs 4 dB) on secondary task performance</i>	100
4.11	<i>Fixed-effect estimates from the LMER model showing the effects of group, RT, SNR on NASA-TLX scores</i>	104
4.12	<i>ANOVA results (Type II Wald χ^2 tests) for NASA-TLX performance</i>	109
4.13	<i>Pairwise group comparisons on NASA-TLX performance across RT and SNR conditions</i>	113
4.14	<i>Pairwise comparison of RT conditions on NASA-TLX performance</i>	121
4.15	<i>Pairwise comparison of SNR conditions (0 dB vs 4 dB) on NASA-TLX performance</i>	126
4.16	<i>Test-retest reliability Cronbach's alpha values for all groups across SNR and RT conditions for primary task performance</i>	131
4.17	<i>Test-retest reliability Cronbach's alpha values for all groups across SNR and RT conditions for secondary task performance</i>	132
4.18	<i>Test-retest reliability Cronbach's alpha values for all groups across SNR and RT conditions for NASA-TLX scores</i>	133

LIST OF FIGURES

Figure No.	Title	Page Number
3.1	<i>Schematic representation of the dual task paradigm</i>	41
3.2	<i>Schematic representation of the procedure to assess the effect of SNR and each RT condition</i>	43
3.3	<i>Schematic representation of the procedure to assess the effect of RT and each SNR condition</i>	44
3.4	<i>Flow chart of procedure adopted in the study</i>	48
4.1	<i>Interaction plots illustrating primary and secondary task performance across signal-to-noise ratio (SNR) conditions, reverberation time (minimum, average, and maximum RT), and speaker groups (Hindi, Malayalam, and Kannada)</i>	57
4.2	<i>Interaction plots illustrating primary task performance across SNR conditions under minimum, average, and maximum RT conditions for Hindi, Malayalam, and Kannada speaker groups</i>	64
4.3	<i>Interaction plots illustrating primary task performance across SNR conditions for the Hindi, Malayalam, and Kannada speaker groups under minimum, average, and maximum RT conditions</i>	69
4.4	<i>Boxplots illustrating primary task performance across RT conditions (minimum, average, and maximum RT) for Hindi, Malayalam, and Kannada speaker groups under 0 dB and 4 dB SNR conditions</i>	72
4.5	<i>Interaction plots illustrating secondary task performance across SNR conditions under minimum and maximum RT conditions for Hindi, Malayalam, and Kannada speaker groups</i>	84
4.6	<i>Interaction plots illustrating secondary task performance across SNR conditions for the Hindi, Malayalam, and Kannada speaker groups under minimum and maximum RT conditions</i>	91
4.7	<i>Boxplots illustrating secondary task performance across RT conditions (minimum and maximum RT) for Hindi, Malayalam, and Kannada speaker groups under 0 dB and 4 dB SNR conditions</i>	96
4.8	<i>Boxplots illustrating NASA-TLX performance across RT conditions (minimum, average, and maximum RT) for Hindi, Malayalam, and Kannada speaker groups under 0 dB and 4 dB SNR conditions</i>	110
4.9	<i>Boxplots illustrating NASA-TLX performance across SNR conditions for the Hindi, Malayalam, and Kannada speaker groups under minimum, average, and maximum RT conditions</i>	116
4.10	<i>Boxplots illustrating NASA-TLX performance across RT conditions (minimum, average, and maximum RT) for Hindi, Malayalam, and Kannada speaker groups under 0 dB and 4 dB SNR conditions</i>	123
4.11	<i>Heat map showing Spearman correlation coefficients between LEAP-Q scores and primary and secondary task performance across different SNR and RT conditions</i>	129
4.12	<i>Heat map showing Spearman correlation coefficients between LEAP-Q scores and NASA-TLX ratings across different SNR and RT conditions</i>	130

ABSTRACT

Speech perception in adverse acoustic environments requires not only accurate auditory processing but also the allocation of cognitive resources, a construct referred to as listening effort. While the effects of background noise and reverberation on speech intelligibility are well-documented, research examining their combined influence on listening effort in multilingual populations like India remains limited. The present study investigated and compared the effects of signal-to-noise ratio (SNR) and reverberation time (RT) on listening effort in native Kannada speakers and two groups of non-native Kannada speakers: those with Hindi (Indo-Aryan) and Malayalam (Dravidian) as their first language.

A total of 90 adults (30 per group) with normal hearing participated in a mixed factorial design. Listening effort was assessed using a dual-task paradigm adapted from the Sentence-Final Word Identification and Recall (SWIR) test, comprising primary (speech recognition) and secondary (word recall) tasks, alongside subjective ratings using the NASA Task Load Index (NASA-TLX). Stimuli were presented across two SNR conditions (0 dB and 4 dB) and three RT conditions (RT minimum = 0.74 s, RT average = 1.66 s, RT maximum = 4.37 s), simulating real-world Indian classroom acoustics. Statistical analyses employed beta generalised linear mixed models (Beta-GLMM) for primary task data, linear mixed-effects regression models (LMER) for secondary task and NASA-TLX data. Spearman correlation was used to examine the relationship between second-language proficiency (LEAP-Q) and all three listening effort measures (primary, secondary, and NASA-TLX).

Results revealed significant main effects of RT, SNR, and language group on both primary and secondary task performance. Increasing RT and decreasing SNR consistently degraded speech recognition and elevated cognitive listening effort. Native Kannada

speakers demonstrated the highest task scores and lowest cognitive load, followed by Malayalam speakers, with Hindi speakers showing the poorest performance. This finding is consistent with the typological proximity of each group's first language to Kannada. Significant two-way and three-way interactions among RT, SNR, and group indicated that non-native speakers were disproportionately affected under combined adverse acoustic conditions. For subjective listening effort, language group was the strongest predictor, while RT did not produce a significant main effect on NASA-TLX scores, revealing an important dissociation between objective and subjective measures of listening effort. Second language proficiency was strongly positively correlated with primary and secondary task performance and moderately negatively correlated with NASA-TLX ratings, with the strongest relationships observed under the most adverse acoustic conditions. Test–retest reliability was good to excellent across all measures (Cronbach's $\alpha = 0.76\text{--}1.00$).

These findings demonstrate that listening effort in adverse acoustic environments is shaped by a complex interplay of acoustic and linguistic factors, with the typological relationship between a speaker's first language and the target language serving as a meaningful predictor of cognitive listening demand. The results have direct implications for audiological assessment, acoustic design of educational environments, and language accessibility in India's linguistically diverse settings.

CHAPTER 1

INTRODUCTION

Speech perception is a complex process beyond recognizing individual sounds and words. It requires the continuous integration of various pieces of information to form a meaningful and coherent understanding of what is being communicated. In challenging acoustic environments, where the speech signal is degraded, relying solely on acoustic cues is often insufficient. In such situations, the speaker's ability to draw on cognitive resources like attention, memory, and language knowledge becomes crucial for successful speech comprehension. Individual differences in cognitive ability have substantial impact on speech recognition when listening conditions are challenging.

Understanding speech in complex and adverse environments can be difficult and effortful, requiring the accurate perception of speech sounds and the continuous integration of linguistic and contextual information to extract meaning. Everyday listening environments such as classrooms, cafés, open-plan offices frequently present adverse conditions including elevated background noise levels (BNL), and extended reverberation times (RT) which degrade signal quality and increase listening difficulty (Mattys et al., 2012). While such adverse conditions are often created by increasing background noise, they can also result from conversational rather than clear speech (van Engen et al., 2012) or accented speech in contrast to speech (Adank et al., 2009). Moreover, speaker characteristics, particularly language proficiency, further modulate listening ease in adverse environments.

Speech perception studies predominantly assess recognition accuracy, often overlooking the underlying cognitive demands required for successful comprehension. This approach limits the practical utility of such tests in real-world communicative settings. In

contrast, listening effort pertains to allocating mental resources during auditory processing, especially in adverse conditions, and may not correlate directly with intelligibility (Pichora-Fuller et al., 2016). This is particularly evident under conditions where speech remains intelligible yet still imposes considerable cognitive load on the speaker. Although speech perception in noise has been extensively examined, particularly regarding intelligibility metrics under varying acoustic conditions, the construct of listening effort remains comparatively underexplored. Strand et al. (2018) investigated how increasing background noise affects listening effort in young adults with normal hearing, and whether this effect is moderated by working memory capacity. Through a dual-task paradigm, they demonstrated that listening effort consistently increased as noise levels rose, while working memory capacity showed no significant relationship with the amount of effort expended.

Researchers have shown that non-native language proficiency significantly improves speech perception in noisy environments and that increased English proficiency enhances performance across various masker types, including both energetic (stationary, fluctuating noise) and informational (two-talker babble) (Kilman et al., 2014). Bilingual speakers exhibited significantly poorer performance on QuickSIN tests than monolinguals, requiring more favourable SNRs to attain the same intelligibility level (Alqattan & Turner, 2021).

Despite growing interest, no consensus exists on the most reliable and sensitive methods to measure listening effort. Various approaches—including physiological (e.g., pupillometry, EEG), behavioural (reaction time, dual-task paradigms), and self-report techniques are used with varying outcomes (Giuliani et al., 2021). Some studies have reviewed listening effort measures in adult populations, and they recommended use of both objective and subjective methods for listening effort measurement (Keur-Huizinga et al., 2024; Tednes & Seeman, 2017). While prior studies have attempted to quantify listening

effort, most have employed a single methodological lens, either objective (e.g., pupillometry, reaction time) or subjective (e.g., self-report ratings) resulting in a limited understanding of its multidimensional nature (Giuliani et al., 2021). Most studies investigating listening effort have employed the dual-task paradigm, where a primary speech perception task is paired with a simultaneous secondary task to gauge cognitive load (Gagné et al., 2017). Common secondary tasks include verbal responses (e.g., lexical decision), manual responses (e.g., button press), and a reaction-time task (Gosselin & Gagné, 2011). Picou et al. (2016) used a semantic dual-task paradigm to assess listening effort. The primary task involved monosyllable word recognition, while the secondary task required participants to identify nouns and press a button as quickly as possible. Longer response times indicated greater listening effort. Their findings showed increased listening effort in noise, but not in reverberation conditions. Similarly, Brown & Strand (2019) demonstrated that listening effort rises with noise level, regardless of working memory capacity.

In addition to behavioural approaches, subjective measures are commonly employed to assess listening effort. The Speech, Spatial and Qualities of Hearing Scale (SSQ) evaluates self-perceived difficulties in understanding speech, spatial hearing, and sound clarity in everyday environments (Gatehouse & Noble, 2004). The Listening Effort Questionnaire- Cochlear Implant (LEQ-CI) (Hughes et al., 2019) and Effort Assessment Scale (EAS) (Alhanbali et al., 2017) specifically target mental effort associated with various listening demands. The NASA Task Load Index (NASA-TLX) has been adapted to assess listening effort. It includes six components: mental demand, physical demand, temporal demand, performance, effort, and frustration (Hart & Staveland, 1988). Visual analog scales, and Likert-type ratings are frequently used post-tasks to capture perceived effort in real time (Pals et al., 2015). These subjective methods, especially with behavioural or physiological

indices, provide a richer understanding of listening effort in diverse and multilingual listening environments.

Listening effort is modulated both by acoustic and intrinsic factors. While SNR and RT are acoustic factors, speaker characteristics, age, hearing status, cognitive capacity, and language proficiency, are intrinsic factors (Desjardins & Doherty, 2014). Using high SNRs, for instance, may yield minimal impact on intelligibility but significantly elevate listening effort, particularly in cases involving synthetic speech or non-native language processing (Smeds et al., 2015). Such effortful listening, especially under adverse acoustic conditions, has been linked to mental fatigue and impaired dual-task performance (Gagné et al., 2017). Likewise, studies on the intrinsic factor of language proficiency consistently show that non-native (proficient in other than the target language or second language - L2) speakers experience greater difficulty and expend more cognitive resources than native (proficient in the target language - L1) speakers, when processing speech in noise or reverberant environments. Even highly proficient non-native speakers demonstrate significantly lower recognition accuracy than native speakers under adverse listening conditions, despite comparable performance in quiet (Takata & Nábělek, 1990). Non-native speakers are disproportionately affected due to reduced familiarity with the phonological and lexical structure of the second language, which further compounds the effort required for successful speech-in-noise perception (Borghini & Hazan, 2018). These difficulties are compounded by unfamiliar accents (Stringer & Iverson, 2020) or reduced semantic context (Borghini & Hazan, 2018), all of which demand greater attentional and working memory resources. Gosselin & Gagné. (2011) found that older adults exert significantly more listening effort than younger speakers, even without measurable hearing loss. Similarly, Picou et al. (2016) reported that while noise and reverberation reduced speech intelligibility, only noise increased cognitive load in young normal-hearing adults.

Previous research in India has explored how language background, masker type, and acoustic conditions influence speech perception. Manikandan & Geetha (2020) investigated the effect of reverberation and noise on speech recognition in normal-hearing listeners and found that increasing reverberation time and decreasing SNR significantly degraded speech scores, with greater difficulty observed in adverse listening conditions. However, this study was limited to native speakers and focused solely on recognition accuracy without assessing cognitive listening effort. Naresh & Yathiraj (2018) compared listening effort between native and non-native speakers under varying SNRs using dual-task paradigms, reporting that non-native speakers required more cognitive resources to maintain speech recognition, especially at lower SNRs. While this addressed the listening effort, reverberation effects were not included, and only one native language was examined. Shastri & Kumar (2015) investigated the role of linguistic distance in speech-in-speech recognition, finding the highest scores with unfamiliar language babble and the lowest with native babble. This work highlighted informational masking effects but did not vary SNRs or reverberation, nor did it assess cognitive effort. Suman and Geetha (2018) studied Kannada–English bilingual’s Kannada sentence recognition in five babble conditions, showing poorest performance with native Kannada babble and best with Nepali babble. The study concluded that masker–target linguistic similarity plays a larger role than language proficiency in masking. However, the research used a single fixed SNR (0 dB) and did not explore reverberation or cognitive effort. The influence of nativity of language is particularly relevant in India which hosts one of the most linguistically diverse populations globally, encompassing hundreds of languages spread across several prominent genealogical families namely Indo-Aryan, Dravidian, Austro-Asiatic, and Tibeto-Burman, as well as smaller families such as Tai-Kadai and Great Andamanese, and language isolates like Nihali and Burushaski (Kulkarni-Joshi, 2019). Despite their disparate origins, these languages have historically coexisted and interacted,

resulting in a linguistic area marked by shared structural features, including retroflex consonants, subject-object-verb word order, and morphological reduplication. This deep multilingual environment, shaped by extensive and sustained language contact, often gives rise to asymmetries in speech perception. Such language shaped experiences particularly become important under adverse acoustic conditions such as background noise or reverberation. In such settings, the differential allocation of cognitive resources by native and non-native speakers becomes especially salient. Given this complex linguistic landscape and the widespread prevalence of bilingualism and multilingualism, investigating listening effort across varying SNRs and reverberant environments provides a crucial window into the cognitive demands of real-world communication in linguistically heterogeneous contexts like India.

1.1 Need for the Study

Much of the literature on listening effort has centred on English as the native language of speakers, often overlooking the linguistic diversity present in multilingual nations like India. In such contexts, the native language background of speakers can introduce considerable variability in phonological, prosodic, and syntactic structures which in turn can reflect in non-native language perception. The cross-linguistic differences can significantly modulate the ease of processing speech cues, thereby influencing the degree of listening effort required (Mattys et al., 2012; Rosenhouse et al., 2006).

While some Indian studies demonstrate that language background, SNR, and masker type affect speech perception (Eranna et al., 2025; Jain et al., 2014; Kanigalpula et al., 2025; Naresh, 2018), they primarily measured accuracy and did not systematically assess listening effort, especially in reverberant conditions. None has jointly examined the combined effects

of SNR variation, reverberation, and speaker nativeness on cognitive listening demands. Given multi-cultural landscape of India in general, and higher noise in real-world communications in particular (Das et al., 2019; Jamir et al., 2014), speakers specifically those non-natives to the language might have to allocate additional cognitive resources to understand speech in noisy and reverberant spaces.

In the Indian context, where multilingualism is the norm rather than the exception, speakers often communicate in a second (L2) or third language across different social and professional settings compared to their first language (L1). While numerous studies have independently explored listening effort in native or non-native populations (Borghini & Hazan, 2020; Brännström et al., 2021; Peng & Wang, 2019; Simantiraki et al., 2023; Song & Iverson, 2018), there is a lack of direct comparative data examining how SNR and reverberation impacts listening effort differently in L1 vs. L2 speakers, particularly using controlled experimental designs. Individuals from different linguistic backgrounds (e.g., Indo-Aryan vs. Dravidian first language) may process a second language like Kannada differently due to phonological, syntactic, and prosodic mismatches, potentially influencing the effort required during listening. There is also a lack of evidence on how the interaction between native language status, SNR, and reverberation shapes listening effort in this population. Moreover, Indian classrooms, workplaces, and public environments often suffer from poor acoustic design with high background noise and reverberation (Saravanan et al., 2019; Sundaravadhanan et al., 2017), exacerbating these challenges. Understanding how these factors influence listening effort in native and non-native speakers can provide critical insights for improving educational accessibility, audiological assessments, and speech-based technologies in India's unique linguistically heterogeneous society.

To address these gaps, the present study examined and compared the listening effort in native and non-native speakers of the Kannada language across varying SNR conditions

and reverberation. By focusing on languages which share similar (e.g., native speakers of Malayalam listening to Kannada) and different (e.g., native speakers of Hindi listening to Kannada) characteristics to Kannada, the study seeks to understand how linguistic familiarity, structural differences, and L2 status interact with acoustic challenges to modulate cognitive load during speech perception.

1.2 Aim of the Study

The present study aimed to investigate and compare the effects of signal-to-noise ratio (SNR) and reverberation time on listening effort in native and non-native Kannada (Kannada-Malayalam; Kannada-Hindi) speakers.

1.3 Objectives of the Study

1. To investigate the effect of SNR (0 and 4 dB) on primary and secondary tasks scores of the dual task paradigm in native (Kannada) and non-native (Kannada-Malayalam; Kannada-Hindi) speakers.
2. To investigate the effect of reverberation time -RT (RT minimum, RT average, RT maximum) on primary and secondary tasks scores of the dual task paradigm in native (Kannada) and non-native (Kannada-Malayalam; Kannada-Hindi) speakers.
3. To investigate the effect of SNR (0 and 4 dB) on subjective listening effort ratings in native (Kannada) and non-native (Kannada-Malayalam; Kannada-Hindi) speakers.
4. To investigate the effect of RT (RT minimum, RT average, RT maximum) on the subjective listening effort ratings in native (Kannada) and non-native (Kannada-Malayalam; Kannada-Hindi) speakers.
5. To explore the relationship between second language proficiency and listening effort

measures of the dual task paradigm or subjective ratings across native and non-native speaker groups, SNRs (0 and 4 dB) and RT (RT minimum, RT average, RT maximum) conditions.

1.4 Null Hypotheses

H1: There is no significant effect of SNR on primary task and secondary task scores of a dual task paradigm in native and non-native Kannada speakers (Kannada-Malayalam; Kannada-Hindi).

H2: There is no significant effect of RT (RT minimum, RT average, RT maximum) on primary task and secondary task scores of a dual task paradigm in native and non-native Kannada speakers (Kannada-Malayalam; Kannada-Hindi).

H3: There is no significant effect of SNR on subjective listening effort ratings in native (Kannada) and non-native (Kannada-Malayalam; Kannada-Hindi) speakers.

H4: There is no significant effect of RT (RT minimum, RT average, RT maximum) on the subjective listening effort ratings in native (Kannada) and non-native (Kannada-Malayalam; Kannada-Hindi) speakers.

H5: There is no significant relationship between second language proficiency and listening effort measures of the dual task paradigm / or subjective ratings across native and non-native speakers across SNRs (0 and 4 dB) and RT (RT minimum, RT average, RT maximum) conditions.

CHAPTER 2

REVIEW OF LITERATURE

The perception of speech is a complex multi-level process that depends on the integration of acoustic-phonetic and linguistic-semantic cues. The effectiveness of integration is controlled by the range of intrinsic properties of the speaker, including features such as language skills, cognitive ability, hearing, and memory, as well as by extrinsic conditions such as presence of ambient noise, reverberation, and other acoustic disturbances (Mattys et al., 2012). Acoustic disturbances have the effect of creating extra mental load, especially for foreign-language speakers.

2.1. Concept and definition of listening effort

The comprehension of speech in natural conditions can involve great mental effort, especially if adverse acoustic conditions such as presence of background noise, reverberation, or interfering speech prevail (Mattys et al., 2012). In this context, the concept of listening effort has gained increasing attention in recent years. Listening effort refers to the allocation of mental effort to auditory perception, especially in difficult acoustic conditions (Pichora-Fuller et al., 2016). Francis & Love (2020) conceptualize listening effort as a multidimensional construct encompassing both cognitive processing demands and affective responses. The cognitive component reflects the allocation of mental resources such as attention and working memory during speech processing, particularly under acoustically adverse conditions, whereas the affective component relates to the speaker's subjective experience, including perceived effort, fatigue, and motivation.

Listening effort has been defined as the deliberate allocation of cognitive resources to overcome obstacles in speech understanding, particularly under challenging conditions (Pichora-Fuller et al., 2016). The Framework for Understanding Effortful Listening (FUEL) proposed by Pichora-Fuller et al. (2016) emphasizes that listening effort is distinct from speech intelligibility and is governed by the interaction of sensory input, cognitive capacity, and task demands. This distinction has led to a growing body of research in both theoretical and clinical domains, with researchers employing a variety of innovative methods to better understand listening effort. Acoustically degraded speech imposes significant cognitive demands. These demands are reflected not only in behavioral performance but also in physiological and neural responses. Listening effort involves the interaction of multiple systems, including auditory processing, working memory, and executive control. As a result, speech perception under challenging conditions is not merely a sensory process but a complex cognitive activity that varies across individuals and contexts (Peelle, 2018).

2.2. Factors affecting listening effort

Listening effort is strongly influenced by the interaction of extrinsic (environmental) and intrinsic (speaker-related) factors, both of which determine how much cognitive resource is required to understand speech, particularly under adverse conditions. Understanding the role of intrinsic and extrinsic factors is crucial because it highlights that speech intelligibility alone is insufficient to capture listening difficulty. Two individuals may achieve similar accuracy but differ significantly in the effort required. This has important implications for clinical assessment, educational settings, and multilingual environments, where increased listening effort can lead to fatigue, reduced comprehension over time, and decreased overall communication efficiency. Therefore, considering both types of factors provides a more comprehensive understanding of listening effort and supports the need for multidimensional assessment and management approaches in research and practice.

2.2.1 Effect of Extrinsic Factors on Listening Effort

Noise and reverberation: Several studies have explored how adverse acoustic conditions influence listening effort. For instance, Picou et al. (2016) used a dual-task paradigm to investigate the effects of noise and reverberation on young adults with normal hearing. While both factors reduced speech recognition accuracy, only background noise significantly increased response times on the secondary task, indicating greater cognitive effort. Interestingly, reverberation, even at higher levels did not appear to increase listening effort, suggesting that noise may place a greater burden on cognitive resources than reverberation under certain conditions.

Similarly, Brown & Strand (2019) demonstrated that listening effort increases consistently as background noise levels rise. However, their findings also indicated that working memory capacity did not significantly influence the degree of effort expended, suggesting that noise alone is a dominant factor in increasing cognitive load.

Recent research examining classroom environments has shown that both children and adults perceive higher listening effort in noisy conditions compared to quiet settings. Using a dual-task approach, Seitz et al. (2024) found that subjective reports of effort were closely aligned with behavioral performance, particularly in school-aged children and adults. This highlights the sensitivity of dual-task paradigms in capturing listening effort across age groups.

Speech rate: Speech rate is another variable that significantly influences listening effort. Faster speech places additional demands on cognitive processing, making it more difficult for speakers to keep up with incoming information. Jeong et al. (2024) showed that

as speech rate increased, participants' performance on a concurrent task declined, reflecting higher cognitive load. Errors in sentence repetition also increased, suggesting that rapid speech disrupts both comprehension and memory processes. Borghini & Hazan (2020) showed that acoustic enhancements, such as clear speech, reduce listening effort more effectively than semantic context, highlighting the importance of signal clarity over contextual cues. These findings emphasize the importance of controlling speech rate to facilitate effective communication.

Accent: Accent variation adds another layer of complexity to speech perception. Accented speech often deviates from familiar phonetic patterns, making it more difficult to process, especially for non-native speakers (Munro & Derwing, 1995). However, speakers can adapt to unfamiliar accents over time. Studies by Adank et al. (2009) and Goslin et al. (2012) show that while both native and non-native speakers are capable of adaptation, the process tends to be slower and more effortful for non-native individuals.

Accent and competing speech further increase this burden. Song & Iverson (2018) demonstrated that non-native speakers require greater neural resources when processing accented or multi-talker speech. Similarly, Peng & Wang (2019) found that for non-native speakers, the effect of accent can be as demanding as noise or reverberation.

Complexity of Task: The complexity of the listening task and environmental conditions (e.g., competing speech, multi-talker scenarios) also influence listening effort. Different speech types also influence listening effort. Simantiraki et al. (2023) reported that while enhanced speech improves intelligibility and reduces effort, synthetic speech tends to increase cognitive load. Interestingly, Lombard speech was associated with lower effort despite not always yielding the highest intelligibility.

Lau et al. (2019) examined how different types of auditory tasks influence physiological and subjective measures of listening effort in individuals with normal hearing. The study found

that listening effort varies significantly depending on the nature and complexity of the task, even when the acoustic conditions remain constant. Physiological measures (such as pupil dilation) and subjective ratings did not always align, indicating that they capture different dimensions of effort. More cognitively demanding tasks elicited greater physiological responses and higher perceived effort. The findings highlight that task type is a critical factor in measuring listening effort and reinforce the need to use multiple assessment methods for a comprehensive understanding.

In addition, difficult listening environments increase reliance on cognitive resources such as attention and working memory, leading to higher effort (Peelle, 2018). Furthermore, recent EEG studies highlight that listening effort is not static but evolves over time. Hunter (2022) demonstrated that neural indicators of effort, such as alpha and theta oscillations, change dynamically with increasing task difficulty and fatigue. This suggests that prolonged listening under adverse conditions leads to cumulative cognitive strain.

2.2.2 Effect of Intrinsic Factors on Listening Effort

In addition to environmental influences, intrinsic speaker characteristics play a significant role in determining listening effort. Factors such as working memory, attention, age, and language proficiency influence how effectively individuals process speech, particularly in degraded conditions.

Cognitive abilities: Individuals with stronger cognitive abilities, especially higher working memory capacity, are better able to compensate for missing or unclear auditory input by relying on contextual and top-down processing strategies (Rönnberg et al., 2013). Electrophysiological studies using EEG have demonstrated that neural markers such as N400 responses are enhanced in non-native speakers, indicating increased semantic processing effort (Romero-Rivas et al., 2015).

Francis et al. (2021) investigated how speaker characteristics influence both self-reported and physiological measures of listening effort under challenging listening conditions. The study found that individual differences, such as language background and cognitive abilities, affect these measures in different ways. While physiological indicators (e.g., pupil dilation) reflected real-time cognitive load, self-reported ratings were more influenced by subjective perception and speaker experience. Importantly, the relationship between these measures was not consistent, suggesting that they capture distinct aspects of listening effort.

Sarampalis et al. (2009) examined the relationship between attention and listening effort during speech processing, particularly under adverse listening conditions. The study demonstrated that when speech is degraded by noise, speakers must allocate greater attentional resources to maintain understanding, resulting in increased listening effort. Using behavioral measures, it was found that even when speech intelligibility remains relatively stable, cognitive load increases significantly, affecting performance on concurrent tasks. The findings highlight that listening effort is closely tied to attentional demands and that successful speech perception in noisy environments requires substantial cognitive resource allocation beyond what is reflected by accuracy alone.

Language experience: Language background plays a critical role in listening effort. Differences in listening effort between native and non-native speakers are particularly evident in noisy environments. Non-native speakers generally require more cognitive resources to process speech, even when their comprehension accuracy is similar to that of native speakers.

Peng & Wang (2019) showed that non-native speakers exhibit higher listening effort compared to native speakers, particularly in environments with noise and reverberation.

This increased effort is attributed to reduced automaticity in language processing and greater reliance on cognitive resources. Similarly, earlier studies (Lecumberri et al., 2010; Takata & Nábělek, 1990) have consistently shown that non-native speakers perform worse under adverse conditions. More recent work using physiological measures has confirmed that they also experience greater listening effort (Borghini & Hazan, 2018).

Non-native speakers, in particular, tend to expend more cognitive effort even when speech intelligibility is comparable to that of native speakers. This increased effort has been demonstrated through both behavioral measures and physiological indicators such as pupil dilation and EEG responses (Borghini & Hazan, 2018).

Brännström et al. (2021) investigated listening effort and fatigue in native and non-native primary school children, focusing on how language background influences cognitive demands during listening tasks. The study found that non-native children experienced significantly higher listening effort and greater fatigue compared to native peers, particularly in challenging listening conditions. Despite achieving comparable levels of speech understanding in some cases, non-native speakers required more cognitive resources, indicating reduced processing efficiency. The findings highlight that language proficiency plays a crucial role in modulating listening effort from an early age and that sustained effortful listening may contribute to increased fatigue over time. This study underscores the importance of considering both effort and fatigue beyond accuracy when evaluating speech perception in multilingual educational settings.

Age: Age is another important factor. McGarrigle et al. (2014) reported that older adults often require greater cognitive resources for speech comprehension, even in the absence of measurable hearing loss. Supporting this, Gagne et al. (2023) found that older adults showed higher listening effort than younger adults under noisy conditions, possibly due to increased reliance on compensatory cognitive strategies. Additionally, Petersen et al.

(2015) demonstrated that both increased memory load and higher noise levels lead to declines in performance and slower responses, indicating elevated effort in older individuals. Similarly, Degeest et al. (2015) found that older adults experience greater listening effort than younger individuals, even when their speech understanding is similar. This is due to age-related declines in auditory and cognitive abilities, leading to increased reliance on cognitive resources during listening.

Hearing Status: Listening effort is also influenced by hearing status. Studies have shown that individuals with hearing loss often experience greater cognitive demands during speech perception, particularly in noisy environments. Nishida et al. (2024) found that under highly challenging conditions, individuals with hearing loss showed reduced accuracy and increased response times, indicating higher effort.

Hicks & Tharpe (2002) reported that children with hearing loss experience greater listening effort and fatigue compared to their normal-hearing peers, especially in classroom listening conditions. This increased effort can lead to reduced attention, poorer academic performance, and higher levels of listening-related fatigue over time.

Similarly, hearing aid users have been reported to experience greater listening effort compared to normal-hearing individuals, despite improvements in audibility (Salah et al., 2024). This suggests that amplification alone may not fully alleviate the cognitive burden associated with hearing impairment.

The scoping review by Philips et al. (2023) reports that cochlear implant users often experience increased listening effort and fatigue, particularly in challenging listening environments, despite improved speech intelligibility. This heightened effort is attributed to the need for greater cognitive resource allocation due to degraded auditory input.

Perreau et al. (2017) investigated listening effort in adults with normal hearing and cochlear implant users using a dual-task paradigm. The study demonstrated that cochlear

implant users exhibited significantly greater listening effort compared to normal-hearing individuals, even when speech recognition performance was relatively similar across groups. This increased effort was reflected in poorer secondary-task performance, indicating a higher allocation of cognitive resources to auditory processing. The findings suggest that degraded auditory input, as experienced by cochlear implant users, imposes additional cognitive demands during speech perception. Importantly, the study highlights that speech intelligibility alone does not fully capture listening difficulty, reinforcing the need to assess listening effort as a complementary measure in both clinical and research settings.

2.3 Assessment of Listening Effort

Assessment of listening effort requires a clear understanding of both intrinsic and extrinsic factors, as these jointly determine the cognitive load experienced during speech perception. Extrinsic factors, such as background noise, signal-to-noise ratio (SNR), reverberation, speech rate, and accent, define the difficulty of the listening environment by degrading the acoustic signal. In contrast, intrinsic factors—including language proficiency, cognitive capacity (e.g., attention and working memory), age, and hearing status determine the speaker's ability to cope with these challenges. The Framework for Understanding Effortful Listening (FUEL) emphasize that listening effort arises from the interaction between task demands and available cognitive resources. Therefore, effective assessment must account for both types of factors rather than treating listening effort as a purely acoustic or purely cognitive phenomenon.

Knowledge of these factors directly informs the selection and interpretation of assessment methods. For example, when extrinsic demands are manipulated (e.g., varying SNR or reverberation), behavioral measures such as dual-task paradigms can capture

changes in cognitive load through reaction time or performance decrements. At the same time, subjective measures (e.g., NASA-TLX, rating scales) reflect the speaker's perceived effort, which is often influenced by intrinsic characteristics such as motivation and fatigue. Physiological measures (e.g., pupillometry, EEG) provide additional objective indices of real-time processing demands. Importantly, because intrinsic factors differ across individuals, the same acoustic condition may yield different levels of listening effort, highlighting the need to interpret results in relation to speaker characteristics such as native versus non-native language status.

Understanding the role of intrinsic and extrinsic factors enhances the clinical and research utility of listening effort assessment. It allows researchers to design ecologically valid experiments that simulate real-world communication, and helps clinicians identify individuals who may experience high effort despite adequate speech intelligibility. This is particularly important in multilingual populations, where non-native speakers often require greater cognitive resources under adverse conditions. By integrating both environmental and speaker-related variables, assessment becomes more comprehensive, enabling better prediction of communication difficulties, listening fatigue, and functional performance in everyday settings.

Listening effort can be measured using three main approaches: behavioral, physiological, and subjective methods. Behavioral measures, particularly dual-task paradigms, are widely used to assess how cognitive resources are allocated during listening tasks. As task difficulty increases, performance on the secondary task typically declines, indicating increased effort. Physiological measures include EEG, pupillometry, and other indicators of neural and autonomic activity, providing objective insights into cognitive load.

Subjective measures, such as questionnaires like the SSQ and NASA-TLX, capture individuals' perceived effort and workload. Although each method has its strengths and limitations, researchers generally agree that a combination of approaches provides the most comprehensive assessment of listening effort (Giuliani et al., 2021).

Alhanbali et al. (2019) demonstrated that listening effort is a multidimensional construct, as different measurement approaches capture distinct aspects of the listening process. Using behavioral (dual-task), physiological (pupillometry), and subjective measures, the study found only weak correlations between these methods, indicating that they do not assess a single unified construct. Behavioral measures reflected performance-related demands, physiological measures indexed real-time cognitive load, and subjective ratings represented perceived effort. Based on these findings, the authors concluded that no single measure can fully capture listening effort, and therefore a combination of methods is necessary for a comprehensive assessment. Picou (2023) provided an overview of current methodologies used to measure listening effort and highlights key challenges in the field, emphasizing that different approaches including behavioral, physiological, and subjective measures often yield inconsistent results because they capture different aspects of effort. It also discussed methodological issues such as lack of standardization, task variability, and difficulty in interpreting results across studies. The author suggested that future research should focus on improving theoretical frameworks, integrating multiple measurement approaches, and developing more ecologically valid tasks to better understand listening effort in real-world conditions.

Klink et al. (2012) reviewed 54 publications and identified three major categories of listening effort assessment methods: physiological reactions, cognitive/behavioral performance, and subjective ratings as shown in Table 2.1.

Table 2.1

Classification of listening effort assessment methods showing physiological, cognitive/behavioral, and subjective rating measures.

Type of Measure	Measure / Method	Study	Major Findings
Subjective Measures	NASA-TLX	Hart & Staveland (1988)	Listening effort and perceived workload increased with task difficulty and degraded listening conditions
	Listening difficulty ratings	Schulte et al. (2008)	Listening difficulty decreased with increasing SNR, even when speech intelligibility approached ceiling levels
	Listening effort ratings	McAuliffe et al. (2012)	Older adults perceived greater listening effort when listening to aged voices despite similar intelligibility scores
	SSQ Questionnaire	Gatehouse & Noble (2004)	Listening effort ratings strongly correlated with hearing impairment and communication disability
	Ease of listening ratings	Greenberg et al. (2003)	Array-processing algorithms improved perceived ease of listening in noise
	Listening comfort ratings	Zakis et al. (2009)	Noise reduction algorithms improved listening comfort and ease of understanding but not intelligibility

	Everyday listening effort questionnaire	Meis Gabriel (2001)	&	Hearing aids significantly reduced perceived listening effort during daily communication
Behavioral / Cognitive Measures	Dual-task paradigm	Downs Crum (1978)	&	Presence of noise significantly increased listening effort, reflected by poorer secondary-task performance
	Dual-task paradigm	Picou et al. (2016)		Noise increased listening effort, whereas reverberation showed limited effects on word categorization response times
	Dual-task paradigm	Sarampalis et al. (2009)		Speech recognition and memory performance worsened as SNR decreased; noise reduction reduced listening effort
	Working memory recall task	Ng et al. (2013)		Recall performance decreased under noisy listening conditions, indicating increased listening effort
	Recall and tracking task	Tun et al. (2009)		Older adults and hearing-impaired speakers demonstrated greater dual-task costs and listening effort
	Reaction time task	Gatehouse & Gordon (1990)		Reaction times improved with hearing aid use even when intelligibility changes were minimal
	Speech recall tasks	Rakerd et al. (1996)		Hearing impairment significantly increased memory-related listening effort

	Working memory tasks	Klatte et al. (2007)	Reverberation and classroom noise significantly impaired memory and speech processing performance
	Speech recall and memory tasks	Surprenant (1999)	Noise reduced syllable recall performance and increased cognitive load
Objective/ Physiological Measures	Pupillometry	Kramer et al. (1997)	Poorer SNR conditions produced larger pupil dilation, reflecting increased listening effort
	Pupillometry	Zekveld et al. (2010)	Pupil dilation decreased as speech intelligibility improved
	Pupillometry	Hyönä et al. (1995)	More difficult language processing tasks produced significantly greater pupil dilation
	Pupillometry	Koelewijn et al. (2012)	Single-talker maskers produced greater listening effort than stationary noise
	Pupillometry	Piquado et al. (2010)	Increased memory load resulted in larger pupil dilation responses
	fMRI	Peelle et al. (2010)	Effortful listening activated superior temporal and inferior parietal cortical regions
	EEG Cortical measures	/ Miles et al. (2017)	Increased cortical alpha activity was associated with greater listening demand

Skin	Mackersie &	Task difficulty significantly increased skin
conductance /	Cones (2011)	conductance and electromyographic
EMG		activity
Heart rate	Mackersie &	Heart rate showed limited sensitivity to
variability	Cones (2011)	listening effort changes

Overall, the existing body of literature clearly demonstrates that listening effort is a complex and multidimensional construct influenced by both acoustic and speaker-related factors. Previous studies have consistently shown that adverse listening conditions such as background noise and reverberation increase cognitive load and negatively affect speech processing, particularly in non-native speakers. Further, behavioural, subjective, and physiological measures have highlighted that listening effort cannot be comprehensively understood using a single assessment method alone.

However, despite substantial research on listening effort, several important research questions remain unanswered. First, previous findings regarding the role of reverberation on listening effort remain inconsistent, with some studies reporting significant effects (Holube et al., 2016; Picou et al., 2016; Puglisi et al., 2021) while others demonstrate minimal influence despite (Cueille et al., 2022; Park et al., 2025) reductions in speech intelligibility. Second, most existing studies have independently examined either noise or reverberation (Gagne et al., 2023; Lee et al., 2022; Park et al., 2025; Seitz et al., 2024), with limited evidence regarding their combined influence on listening effort. Third, although non-native speakers are known to experience greater listening effort than native speakers, very few studies have systematically compared native and non-native speakers across varying SNR and reverberation conditions using both behavioural and subjective measures

simultaneously. Most importantly, research examining listening effort in multilingual Indian populations remains scarce, particularly among Kannada native and non-native speakers from different linguistic backgrounds such as Indo-Aryan (Hindi) and Dravidian (Malayalam) language groups.

The present study seeks to address these gaps by systematically investigating the effects of signal-to-noise ratio and reverberation on listening effort in native and non-native Kannada speakers using both behavioural (dual-task paradigm) and subjective (NASA-TLX) measures. Specifically, the study compares native Kannada speakers with non-native Kannada speakers having Hindi and Malayalam as their first languages to understand how linguistic background and second-language proficiency interact with adverse acoustic conditions to influence listening effort. The findings of the study are expected to contribute toward a more ecologically valid understanding of speech perception in multilingual Indian contexts and may have important implications for audiological assessment, classroom acoustics, rehabilitation strategies, and speech-based technologies.

CHAPTER 3

METHODS

3.1 Research Design

The study adopted a mixed factorial design incorporating within-subject and between-subject factors (Orlikoff et al., 2015). The between-subjects factor was the language-based speaker group classification, comprising of native and non-native Kannada speakers. The within-subject factors were SNR and RT, with all participants tested across 2 SNR levels (0 and 4 dB) and 3 RT (RTmin, RTavg, RTmax) conditions.

3.2 Participants

A purposive sampling technique was used to recruit a total of 90 adults belonging to three speaker groups. 30 participants were considered in each group. The age range of participants in all groups varied between 18 and 40 years. The three groups considered in the study are:

- 1) Non- Native Kannada speakers with first language Hindi (Indo-Aryan) (Mean age: 24.70 ± 1.69 y SD)
- 2) Non- Native Kannada speakers with first language Malayalam (South-Dravidian), (Mean age: 23.13 ± 1.57 y SD)
- 3) Native speakers of Kannada with first language Kannada; (Mean age: 21.33 ± 2.19 y SD)

Sample size was calculated using G*Power software Faul et al. (2007). Based on the study by Peng & Wang (2019), for the effect size of 0.35 and power of 0.8, the total sample

size needed for three groups was 82. To account for attribution and variability, a sample size of 90 (30 in each group) was considered.

Informed consent and ethical clearance. Informed consent was obtained from all the participants at the start of the study. A sample of informed consent form is provided in Appendix A. The study was approved by AIISH Institute Review Board (IRB) (Ref no: SH/IRB/M.4/AUD.22/2025-26; Dated: 06/01/2026), as shown in Appendix B.

Inclusion and exclusion criteria: The following inclusion and exclusion criteria were followed while recruiting the participants for the study.

Inclusion criteria:

- a. Adults with normal hearing sensitivity in both ears with their air conduction and bone conduction thresholds within 15 dB HL (Goodman,1965) and air-bone gap not more than 10 dB HL.
- b. Bilateral ‘A’ or ‘As’ type tympanograms with present ipsilateral and contralateral reflexes.
- c. Speech identification scores $\geq 90\%$ in both ears.
- d. All participants must pass the Screening checklist for auditory processing in adults. (SCAP-A) (Vaidyanath & Yathiraj, 2014) as shown in Appendix C.
- e. Normal cognitive abilities as evaluated using Neuropsychological Evaluation Screening Tool (NEST) (Chopra et al., 2018) as shown in Appendix D.
- f. Native Kannada speakers with Kannada as their first language/mother tongue (L1). The participants were exposed to Kannada from early childhood and were fluent users of the language.

- g. Non-native speakers with Hindi or Malayalam as L1 and at least one year of Kannada exposure as second language (L2).
- h. All groups must score above the cut-off criterion on the Modified Language Proficiency Questionnaire.(Yathiraj et al., 2018)
- i. No reports of any otological, neurological, speech and language problems.

Exclusion criteria:

- a. Presence of hearing loss or abnormal middle ear findings.
- b. Speech identification scores below 90%.
- c. Less than one year of Kannada exposure in non-native speakers.
- d. Scores below cut-off on the Modified Language Proficiency Questionnaire (LEAP-Q).
- e. Presence of auditory processing deficits on SCAP-A.
- f. Presence of cognitive impairment on NEST.
- g. Presence of otological, neurological, speech, or language disorders or presence of any illness/health issues during testing.

Table 3.1. shows the demographic, audiological, cognitive, and language characteristics of the participants of the study in three groups.

Table 3.1

Demographic, audiological, cognitive, and language characteristics of the participants

Participant Characteristics	Group 1	Group 2	Group 3
	Mean $\pm$ SD	Mean $\pm$ SD	Mean $\pm$ SD
Age (in years)	24.70 $\pm$ 1.69	23.13 $\pm$ 1.57	21.33 $\pm$ 2.19
Gender	M= 5, F=25	M=2, F= 28	M= 5, F=25

PTA of right ear in dB HL (average thresholds at 0.5kHz, 1 kHz, 2 kHz, 4 kHz)	8.42 ± 2.30	8.58 ± 2.22	8.83 ± 2.13
PTA of left ear in dB HL (average thresholds at 0.5kHz, 1 kHz, 2 kHz, 4 kHz)	9.17 ± 2.21	9.17 ± 2.49	9.46 ± 2.19
Tympanometry (B/L)	A Type=27 As Type=3	A Type=30 As Type=0	A Type=28 As Type=2
SIS in %	100.00 ± 0.00	100.00 ± 0.00	100.00 ± 0.00
LEAP-Q	31.83 ± 0.83	34.53 ± 0.57	58.90 ± 0.80
SCAP-A	1.23 ± 0.43	1.00 ± 0.00	1.00 ± 0.00
NEST	0.53 ± 0.51	0.57 ± 0.50	0.63 ± 0.49

Note. PTA = pure tone average; dB HL = decibels hearing level; SIS = speech identification scores; LEAP-Q = Language Experience and Proficiency Questionnaire; SCAP-A = Screening Checklist for Auditory Processing in Adults; NEST = Neuropsychological Evaluation Screening Tool; SD = standard deviation; M = male; F = female

3.3 Preliminary assessment for participant inclusion:

A detailed preliminary assessment was carried out for all participants prior to data collection to ensure that they met the inclusion criteria for the study. Audiological evaluation included pure tone audiometry, immittance evaluation, and speech identification testing.

Test for screening normal hearing sensitivity - Pure-tone audiometry. Pure tone audiometry was carried out using a calibrated clinical audiometer (Inventis Piano Audiometer / GSI 61 Audiometer) with TDH-39 headphones and a B-71 bone vibrator. Air conduction thresholds were obtained at octave frequencies from 0.25 to 8 kHz, and bone

conduction thresholds were obtained from 0.25 to 4 kHz. Pure tone average (PTA) was calculated by averaging the thresholds (in dB HL) obtained at 0.5 kHz, 1 kHz, 2 kHz, and 4 kHz for both ears using the modified Hughson–Westlake procedure (Carhart & Jerger, 1959). Participants were considered to have normal hearing sensitivity when air conduction and bone conduction thresholds were within 15 dB HL (Goodman, 1965), with an air–bone gap not exceeding 10 dB HL.

Test for screening for speech perception abilities in quiet - Speech audiometry.

Speech audiometry was carried out in a sound-treated room using standardized recorded speech material presented through TDH-39 headphones connected to a calibrated clinical audiometer. Speech recognition thresholds were obtained using ‘The Common Speech Discrimination Tests for Indians’ (Mayadevi, 1974) in Kannada language. Speech identification scores were obtained using ‘Phonetically Balanced Word List’. Most comfortable loudness level measurement for speech was also done for all the clients.

Speech Identification Scores (SIS) were obtained monaurally for both ears at a comfortable listening level 40 dB above the speech recognition threshold. Participants were instructed to listen carefully and repeat the words presented to them. The responses were scored based on the number of correctly identified words, and the percentage score was calculated for each ear. Participants obtaining speech identification scores of 90% or higher in both ears were considered to have normal speech perception abilities and were included in the study.

Test for middle ear status - Immittance evaluation. Middle ear function was tested using a calibrated Immittance meter (GSI Tympanometer Pro). Tympanogram was obtained using a standard 256 Hz probe tone. The ipsilateral and contralateral acoustic reflexes were obtained at 0.5 kHz, 1 kHz, 2 kHz and 4 kHz. Participants exhibiting bilateral ‘A’ or ‘As’

type tympanograms with the presence of both ipsilateral and contralateral acoustic reflexes were included in the study.

Test for evaluating proficiency. Language background and proficiency were also assessed during the preliminary evaluation. Native Kannada speakers had been exposed to Kannada from early childhood and spoke the language fluently. Non-native Kannada speakers had Hindi or Malayalam as their first language and were sequential bilinguals with at least one year of experience using Kannada as their second language. Kannada language proficiency was assessed using the Language Experience and Proficiency Questionnaire (LEAP-Q), adapted to the Indian context by Maitreyee & Goswami (2009) and modified by Yathiraj et al. (2018). The questionnaire is given in Appendix E.

The modified LEAP-Q consisted of eight self-administered questions and the questionnaire was further modified by removing the reading and writing components. The modified questionnaire included sections assessing duration of Kannada use, self-rated proficiency in understanding and speaking, language-switching ability, frequency of language use in everyday situations, native-like accent identification, and daily language exposure. Participants were instructed to respond to all items by selecting the most appropriate response corresponding to their Kannada language use and proficiency. The questionnaire primarily assessed the auditory-verbal domains of understanding and speaking. Responses were scored on a 4-point rating scale, with higher scores indicating greater proficiency and language exposure. The duration-of-use and self-rated proficiency sections contributed maximum scores of 8 each, language-switching ability contributed 4, frequency of language use contributed 56, and native-like accent identification and daily language exposure contributed 4 each. Thus, the total maximum obtainable score for the modified questionnaire was 84, with higher total scores reflecting greater Kannada language proficiency and functional language use. Native Kannada speakers who obtained a

minimum score of 50 and non-native Kannada speakers who obtained a minimum score of 30 were included in the study.

3.4 Procedure

The study was conducted in 3 phases.

Phase 1: Development of test module and software adaptation for testing
(application running on paradigm).

Phase 2: Pilot Testing

Phase 3: Administration of listening effort test

3.4.1. PHASE 1: Development of dual task paradigm test module and software adaptation for testing (application running on paradigm)

The Sentence-Final Word Identification and Recall (SWIR) test developed by Ng et al. (2013) and adapted to Kannada by Shetty et al.(2018) was modified to evaluate listening effort of participants through dual task experiment for the present study. In the first phase, the listening effort test module was modified, with adaptations related to stimulus generation and presentation. The equipment/ softwares involved in the adaptation process is shown in Table 3.2.

Table 3.2

Equipment/ software involved in the adaptations related to stimulus generation and presentation of dual task paradigm.

Equipment/ Softwares	Company specifications	and Purpose
MATLAB Version (R2024b)	Mathworks Inc	For stimulus generation at specified RT and SNR
Paradigm Player Software	Paradigm Inc	Used to run the dual-task paradigm for measuring listening effort.
Laptop	Dell Inspiron with Intel Core i3 Processor	Used for stimulus presentation and running the experimental program.
Transducer	Sennheiser HDA	Used for delivering auditory stimuli during the listening effort task.

Adaptions related to test stimuli. Standardized Kannada sentence lists developed by Geetha et al. (2014) were used as the target stimuli in the dual-task paradigm to assess listening effort. From the Kannada sentence lists, 13 lists (out of 25 lists in the test material) were randomly selected. From these, a total of 130 sentences were obtained, out of which 125 sentences were utilized to assess listening effort across two different SNRs and three RT conditions. The sentences were independent of contextual cues, and the final words were designed to be difficult to predict. Six multi-talker babble noise consisting of three male and three female talkers from the standardized Kannada sentences was mixed with the target sentences. The rationale for using six multi-talker babble noise was to replicate energetic masking and prevent informational masking effects that could interfere with sentence recall.

The stimuli were generated at two SNR conditions (0 dB and +4 dB) and three RT conditions (RT minimum = 0.74 s, RT average = 1.66 s, and RT maximum = 4.37 s) to assess the participants' recognition and retention abilities. The selected SNRs included +4 dB, representing favourable listening conditions similar to a quiet classroom or workplace with mild background noise, where intelligibility is typically near ceiling for native speakers but non-native speakers may still experience measurable listening effort. The 0 dB SNR represented moderate to adverse listening situations, such as group discussions or public places where target and masker levels are equal, conditions known to substantially increase listening effort in non-native speakers (Aparna & Hemanth, 2019; Hiremath & Hemanth, 2019; Shetty et al., 2023). The use of these SNRs (0 and +4 dB) was supported by previous Indian studies (Aparna & Hemanth, 2019; Hiremath & Hemanth, 2019; Janardhan & Devi, 2012; Shetty et al., 2023) as well as related international literature (McGarrigle et al., 2014; Ng et al., 2013; Peng & Wang, 2019).

The selected reverberation conditions (RT minimum = 0.74 s, RT average = 1.66 s, and RT maximum = 4.37 s) were designed to simulate typical graduate-level classroom dimensions in the Indian context (approximately 11 m $\times$ 8.28 m $\times$ 3.6 m) as described by Abraham & Ravishankar (2021). Three RT conditions representing different acoustic environments were replicated: (i) RT minimum (0.74 s), representing a classroom with chairs occupied by students, tables, a 10 mm soft carpet on the floor, and calico cloth-lined curtains; (ii) RT average (1.66 s), representing an empty classroom with chairs, tables, and calico cloth-lined curtains; and (iii) RT maximum (4.37 s), representing an empty classroom without curtains, chairs, or carpet.

The functionality of the MATLAB code developed by Kanigalpula et al. (2025) as modified for classroom dimensions suitable for adult speakers with the input of SNR and RT parameters. The program contained an executable graphical user interface (GUI) that

facilitated the inclusion of an SNR adjustment feature, where signal-in-noise for two SNRs (0 and +4 dB) and three RT conditions (RT minimum = 0.74 s, RT average = 1.66 s, and RT maximum = 4.37 s) were simulated and fed as inputs into the MATLAB code.

The simulation was based on “Room Impulse Response Simulation with the Image-Source Method and Head Related Transfer Function (HRTF) Interpolation,” an inbuilt MATLAB function. The convolution for spatial separation was performed using the “ARI HRTF Database.” The image-source method, described by Allen & Berkley (1979), was used to model specular reflections between the source and receiver. The stimuli convolved with the specified SNR and reverberation conditions were stored as ‘.wav’ output files, after execution of the MATLAB code.

Adaptions related to stimulus presentation. The generated stimuli were then loaded into the Paradigm Experiment Builder and Experiment Player (Version 2.5.0.68, Perception Research Systems Incorporated, Lawrence, Kansas, United States of America) for administration of the dual-task paradigm. A code developed by Spandita (2025) was modified for the present study. Using the Paradigm software, both the primary task involving repetition of the last word of the sentence and the secondary task involving recall of the last but one final words of all sentences within a block were generated and programmed for stimulus presentation and response acquisition.

Irrespective of the SNR or RT condition, the inter-stimulus interval within a block was set at 3000 milliseconds. The inter-block interval (IBI), i.e., the duration between the presentation of two consecutive blocks, was set at 10,000 milliseconds. A beep sound was programmed to be heard after the completion of 5 sentences in a block, following which they had to recall and report the last but one word of the five sentences heard within the block. Five blocks were framed from 25 sentences across each SNR and RT condition to be

tested. A total of 25 blocks [$5 \text{ (blocks)} \times 2 \text{ (SNRs)} + 5 \text{ (blocks)} \times 3 \text{ (RTs)}$] were generated for the study.

To obtain a subjective measure of listening effort and perceived workload during the experimental task, the NASA Task Load Index (NASA-TLX) developed by Hart & Staveland (1988) was modified for the current study. The NASA-TLX is a multidimensional rating procedure that provides an overall workload score based on ratings across six subscales. These domains included mental demand (“How mentally demanding was the task?”), physical demand (“How physically demanding was the task?”), temporal demand (“How hurried or rushed was the pace of the task?”), effort (“How hard did you have to work to accomplish your level of performance?”), frustration (“How insecure, discouraged, irritated, stressed, or annoyed were you?”), and perceived performance (“How successful were you in accomplishing the task?”). Soon after the completion of each block, the software was programmed to display questions from NASA TLX, and the participants ratings on all domains of NASA TLX were programmed to be stored on an excel output. Appendix F shows the NASA-TLX assessment tool.

3.4.2. PHASE 2: Pilot testing

In the second phase, pilot testing was carried out on five native Kannada-speaking adults to evaluate the feasibility, clarity, and administration procedure of the developed dual-task listening effort paradigm. The pilot testing helped verify the appropriateness of the stimuli, software functioning, block arrangement, timing of stimulus presentation, and response recording procedures. Participants with audiological and language expertise similar to those in Group 3 (Native Kannada speakers, see Table 3.1) were administered the listening effort task across the different SNR and RT conditions using the Paradigm software.

During the pilot phase, it was observed that the original task involving repetition and recall of the last word of the sentence posed difficulty for the participants. In Kannada sentences, the last word is predominantly a verb, which was found to be comparatively difficult to perceive, repeat, and recall in noisy and reverberant conditions. Based on these observations, modifications were made to the task structure. Instead of repetition and recall of the last word, the task was changed to repetition and recall of the second-last word or the last but one word of the sentence. The second-last word was typically a noun, which was easier for participants to perceive, interpret, repeat, and recall accurately. This modification improved task clarity and participant performance and was subsequently incorporated into the final experimental protocol.

An example of the modification made during pilot testing is provided below:

Kannada Sentence:

ಅವನು ಮಾರುಕಟ್ಟೆಯಿಂದ ಹಣ್ಣುಗಳನ್ನು ತಂದನು

(Avanu maarukattēyinda hannuḡaḷannu tandanu)

Last word (verb): ತಂದನು (tandanu – brought)

Second-last word / last but one word (noun): ಹಣ್ಣುಗಳನ್ನು (hannuḡaḷannu – fruits)

During pilot testing, participants found it difficult to accurately perceive and recall the sentence-final verbs such as “ತಂದನು” (/tandānu/; *brought*), “ಬೆಳೆದನು” (/beḷeḍānu/; *grew/raised*), and “ತಯಾರಿಸಿದರು” (/təja:risaiḍaru/; *prepared*) etc under noisy and reverberant conditions in the secondary task. However, the second-last words such as “ಹಣ್ಣುಗಳನ್ನು” (/han:uḡaḷannu/; *fruits*), “ತರಕಾರಿಗಳನ್ನು” (/taraka:riḡaḷannu/; *vegetables*), and “ಊಟವನ್ನು” (/u:ṭavannu/; *meal/food*) which are noun with greater semantic salience, were perceived and recalled more consistently. Therefore, the primary

and secondary tasks were modified to include repetition and recall of the second-last word instead of the final word.

The overall accuracy scores of the participants after incorporating the last but one modification in the dual task paradigm ranged between approximately 80% and 90% for both the primary and secondary tasks as shown in the Table 3.3. This indicated that the task difficulty level was appropriate and that the participants were able to understand and perform the task effectively.

Table 3.3

Pilot Testing Performance Scores for Primary and Secondary Tasks Across Participants

Participant	Primary Task Score	Primary Task Mean $\pm$ SD	Primary Task (%)	Secondary Task Score	Secondary Task Mean $\pm$ SD	Secondary Task (%)
1	114	4.56 $\pm$ 0.50	91.2	101	4.04 $\pm$ 0.81	80.8
2	110	4.40 $\pm$ 0.66	88.0	104	4.16 $\pm$ 0.74	83.2
3	116	4.64 $\pm$ 0.48	92.8	98	3.92 $\pm$ 0.89	78.4
4	112	4.48 $\pm$ 0.58	89.6	106	4.24 $\pm$ 0.69	84.8
5	113	4.52 $\pm$ 0.51	90.4	102	4.08 $\pm$ 0.80	81.6

3.4.3. PHASE 3: Administration of the listening effort test on adults

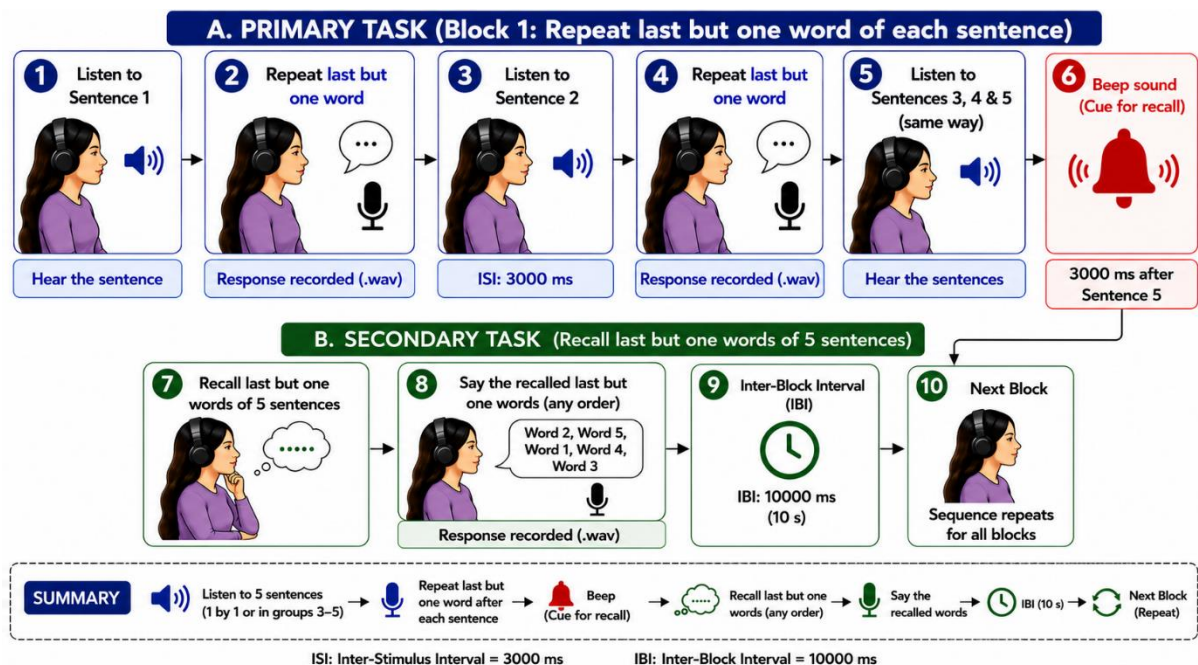
In the third phase, the listening effort test was administered on participants from three groups: native Kannada speakers, non-native Kannada speakers with Hindi as their first language, and non-native Kannada speakers with Malayalam as their first language. To

measure the allocation of cognitive resources, it was proposed to use two testing paradigms: a behavioural measure (dual-task paradigm) and a subjective measure (NASA TLX ratings).

Dual task paradigm. The dual-task paradigm comprised primary and secondary tasks. In the primary task, the speakers were instructed to repeat the second-last word or last but one word of each sentence. In the secondary task, the speakers were instructed to remember the second-last words of the sentences for later recall after completion of each block. The modification from the last word to the second-last word was made following pilot testing (see section 3.4.2). A schematic representation of the dual-task paradigm is presented in Figure 3.1.

Figure 3.1

Schematic representation of the dual task paradigm



The present study employed two SNRs (0 and +4 dB) and three RT conditions conditions (RT minimum = 0.74 s, RT average = 1.66 s, and RT maximum = 4.37 s) to

capture a range of acoustic challenging listening environments. This combination of SNR and RT values enabled investigation of both the independent and interactive effects of noise and reverberation on listening effort across native and non-native Kannada speakers.

Presentation of Stimuli

The stimuli generated using the MATLAB code (see section 3.4.1) were loaded into the Paradigm Experiment Builder and Experiment Player (Version 2.5.0.68, Perception Research Systems Incorporated, Lawrence, Kansas, United States of America) for administration of the dual-task paradigm. A code running on the Paradigm software developed by Spandita (2025) was modified for the present study. Both the primary task involving repetition of the second-last word or last but one word of the sentence and the secondary task involving recall of the second-last words of all sentences within a block were generated using the Paradigm software.

To study the effect of SNR at each RT condition (objective 1; see section-1.3) 25 Kannada sentences along with six multi-talker babble noise were presented randomly at +4 dB and 0 dB SNR. The sentences were presented at the participants' most comfortable level (MCL) for all three groups. The MCL was determined by obtaining the average value across two trials using diagnostic audiometer (Inventis Piano).

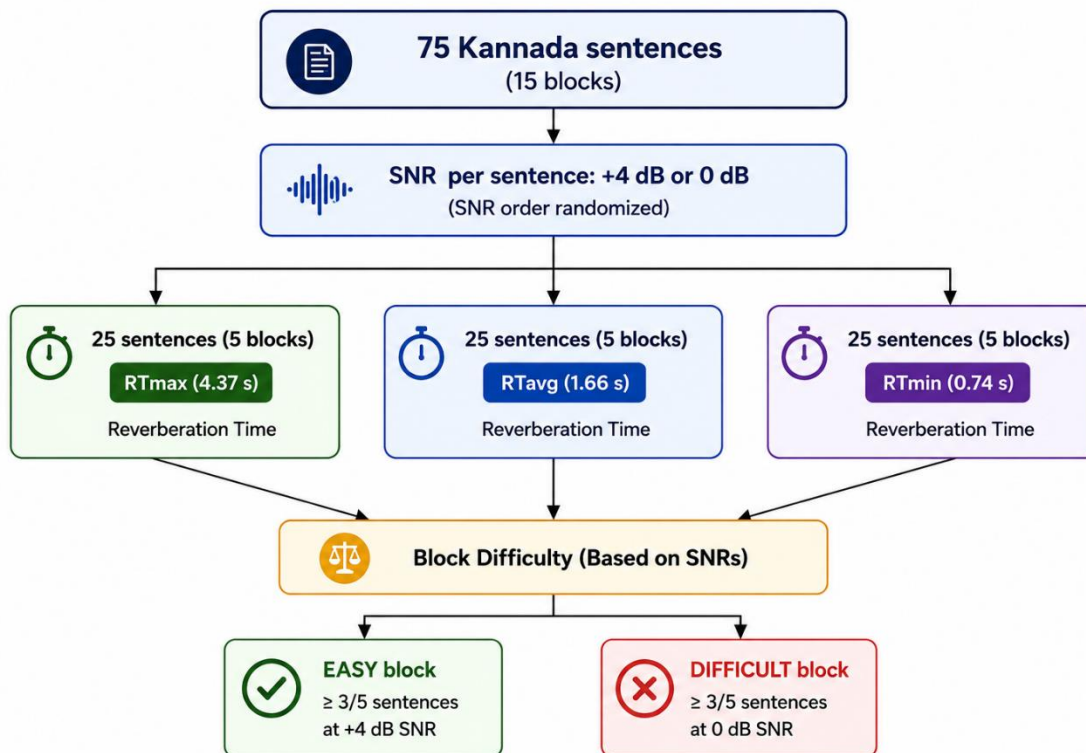
There were five sentences in each block. To ensure balanced exposure, each of the 25 sentences was reverberated in one of the three RT conditions. Thus, five blocks were framed from 25 sentences embedded in noise across each RT condition, constituting a total of 15 blocks (75 sentences) representing all three RT conditions, namely RTmax (4.37 s),

RTavg (1.66 s), and RTmin (0.74 s). In each block for the effect of SNR, a minimum of two instances of each SNR was included.

To account for task difficulty, each block was classified as either “easy” or “difficult” depending on its SNR composition. A block consisting of five sentences was considered “easy” if it contained a majority of sentences presented at +4 dB SNR (three or more sentences), whereas it was considered “difficult” if it contained a majority of sentences presented at 0 dB SNR (three or more sentences). Figure 3.2 illustrates the distribution of sentences in each block for the measurement of the effect of SNR.

Figure 3.2

Schematic representation of the procedure to assess the effect of SNR and each RT condition.

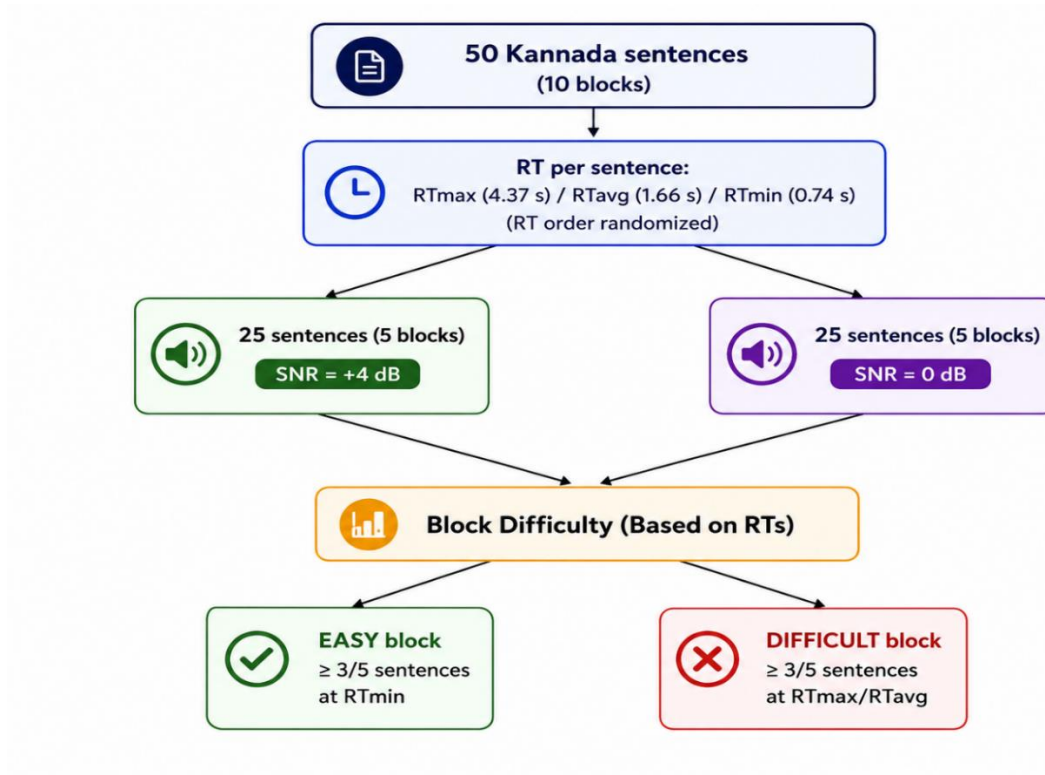


To study the effect of RT at each SNR condition (objective 2, see section-1.3) 25 Kannada sentences along with six multi-talker babble noise were presented randomly across three RT conditions, namely RT minimum (0.74 s), RT average (1.66 s), and RT maximum (4.37 s). The sentences were presented at the participants' most comfortable level for all the groups. There were five sentences in each block. Each of the 25 sentences was paired with one of the two SNR conditions to ensure balanced exposure. For example, one set of 25 sentences mixed with noise was presented at +4 dB SNR, and a similar procedure was followed for another set of 25 sentences mixed with noise at 0 dB SNR. Thus, five blocks were framed from 25 sentences across each SNR condition, constituting a total of 10 blocks (50 sentences) representing the two SNR conditions (+4 dB and 0 dB).

Each block was classified as either “easy” or “difficult” to account for task difficulty, depending on its RT composition. Within each SNR condition (+4 dB and 0 dB), block difficulty was determined based on the RT composition of the five sentences in each block (RT minimum = 0.74 s, RT average = 1.66 s, and RT maximum = 4.37 s). Blocks were classified as “easy” when three or more sentences ($\geq 3/5$) were presented at RT minimum, whereas blocks were classified as “difficult” when three or more sentences ($\geq 3/5$) were presented at RT maximum or RT average. Figure 3.3 illustrates the distribution of sentences in each block for the measurement of the effect of RT.

Figure 3.3

Schematic representation of the procedure to assess the effect of RT and each SNR condition.



The order of presentation of the blocks was randomized and counterbalanced across participants to minimize order effects. The estimated time required for completing the primary and secondary tasks within the dual-task paradigm, along with the subjective listening effort rating for each block, was approximately 2 minutes per block. With a total of 25 blocks, the overall testing duration was approximately 50 minutes. A 10-minute break was provided after every 25 minutes of testing to minimize fatigue and maintain consistent participant performance throughout the experiment.

Subjective rating measure. Immediately after completion of each block in the dual-task paradigm, the participants were asked to complete the NASA-TLX rating. The participants were informed that the NASA-TLX was a subjective assessment and were encouraged to interpret the prompts according to their own perception and experience. Each participant completed 25 sets of NASA-TLX ratings corresponding to the different SNR and RT conditions tested in the study.

Tasks Instructions and Response Acquisition

The stimuli loaded into the presentation software were delivered diotically through headphones at the speakers' most comfortable level, approximately 65 dB SPL. The listening task comprised both primary and secondary tasks. In the primary task, the participants were instructed to repeat the second-last word or last but one word of the heard sentence. In the secondary task, the repeated last but one word of the five sentences within a block were required to be recalled in free order immediately after a beep sound was presented through the headphones.

The speakers were seated in a quiet room during testing. They were instructed to ignore the background noise and listen carefully to the entire sentence. Immediately after hearing each sentence, they were asked to repeat the last but one word of the sentence. The responses were recorded and stored as '.wav' files for final analysis.

Before the commencement of the actual testing, the participants were provided with a practice trial for familiarization purposes. Ten sentences presented as two blocks were administered at +8 dB SNR under both short reverberation (RT minimum) and long reverberation (RT maximum) conditions, with five sentences in each block. The participants were instructed to repeat the last but one word of each sentence and subsequently recall them after hearing the beep sound. The practice items were not included in the final scoring. Participants were encouraged to guess whenever they were uncertain of the response.

Scoring

In the primary task, the response was considered correct when the repeated last but one word was the same as the presented word. Repetition of the last but one word was awarded a score of 1 for a correct response and 0 for an incorrect response or no response. Five sentences were included in each of the 25 blocks, accounting for a maximum score of 125 across the SNR and RT conditions. The following formula was used to calculate the repeat score.

Primary Task – Repeat Score

$$\text{Raw score} = \sum \text{Sum of scores in each block}$$

In the secondary task, the responses were considered correct when the recalled words matched those reported previously in the primary task. No scores were provided when there was no response or when the words were incorrectly recalled. A score of 1 was awarded for each correctly recalled last but one word of the sentence. A maximum score of 5 was awarded when all the last but one words from the five sentences within a block were recalled in any order. The following formula was used to calculate the recall score.

Secondary Task – Recall Score

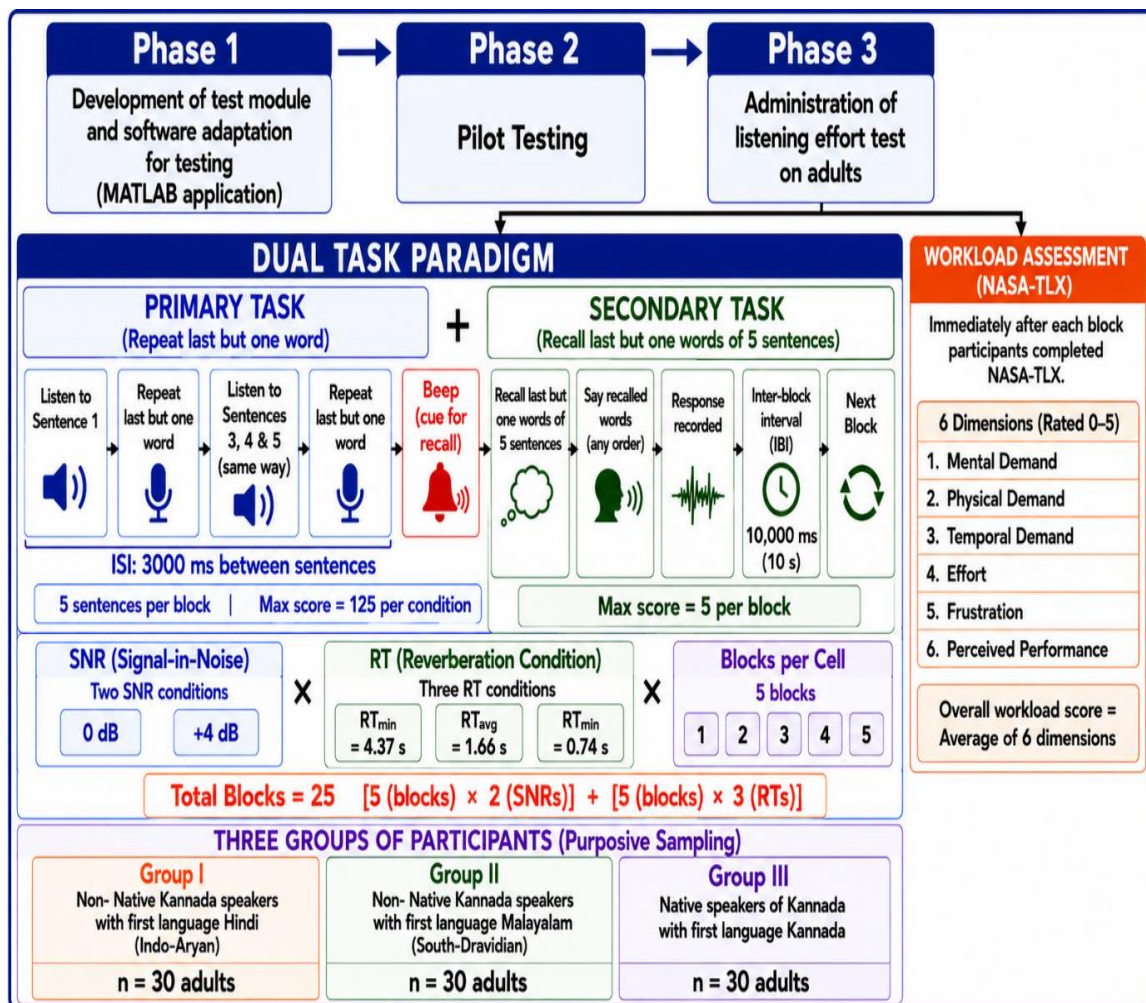
$$\text{Raw score} = \sum \text{Recall score (Recall count of each block)}$$

For the NASA-TLX, although the original NASA-TLX included a pairwise comparison procedure to rank the six subscales and assign weighting scores, this procedure was not performed in the present study. The pairwise comparison was excluded to reduce the overall duration of the experiment and minimize participant-related variability. This

modification was supported by later recommendations provided by Hart (2006). In addition, the rating scale was modified for the present study. Participants rated each item on a scale ranging from 0 (very low) to 5 (very high) using a continuous slider presented through a digital interface. The overall workload score was calculated as the average of the ratings obtained across the six dimensions. Figure 3.4 is the flow chart of procedure adopted in the study.

Figure 3.4

Flow chart of procedure adopted in the study



3.5 Test-retest reliability

The test-retest reliability was determined for 10% of the overall sample size. A total of 9 adults, 3 in each group were re-administered using dual task paradigm after a gap of one month.

3.6 Statistical Analysis

The statistical analyses employed in the present study were selected based on the repeated-measures experimental design and the distributional characteristics of the dependent variables. Since the study involved repeated observations obtained from the same participants across different RT (RT minimum, RT average, RT maximum) and SNR (0 and + 4 dB) conditions, mixed-effects modeling approaches were considered most appropriate. Mixed-effects models are widely recommended in auditory and speech perception research because they simultaneously account for fixed effects associated with experimental manipulations and random effects arising from inter-individual variability (Gafoor & Uppunda, 2023; Koerner & Zhang, 2017; Muhammad, 2023). In addition, these models allow the examination of both main effects and interaction effects among variables such as RT, SNR, and speaker group, thereby providing a comprehensive understanding of how acoustic and linguistic factors jointly influence performance outcomes.

Primary Task

Primary task performance was analyzed using a beta mixed-effects generalized linear model (Beta-GLMM) implemented in RStudio using the glmmTMB package (Brooks et al., 2017). Beta-GLMM was employed because the response variable consisted of proportion-based bounded scores transformed into scaled response scores (SRS) ranging between 0 and 1. Primary task scores were transformed into SRS and analyzed using a beta-

family distribution with a logit link function. Traditional linear regression models are often unsuitable for bounded proportion data because they may violate assumptions of normality and homogeneity of variance (He & Lee, 2022; Maluf et al., 2022). Beta regression models provide a flexible framework for analyzing continuous proportion data while appropriately modeling the heteroscedasticity commonly observed in such measures.

The inclusion of mixed-effects structure further enabled the modeling of repeated observations within participants and facilitated the examination of the fixed effects of Reverberation Time (RT minimum = 0.74 s, RT average = 1.66 s, and RT maximum = 4.37 s), Signal-to-Noise Ratio (SNR: 0 dB and 4 dB), and Group (Hindi, Malayalam, and Kannada) on speech recognition performance during the dual-task paradigm, apart from examining their interactions. Participant ID was included as a random intercept to account for inter-individual variability across repeated observations. Similar generalized linear mixed-effects approaches have been widely used in audiology and speech perception research involving speech recognition scores and perceptual outcome measures (Kim et al., 2025; Xu et al., 2013; Yu et al., 2022).

Model diagnostics were carried out using the Diagnostics for Hierarchical Regression Models using Randomized Quantile Residuals (DHARMA) package (Hartig, 2024) in Rstudio to assess whether the assumptions of the mixed-effects models were adequately satisfied. DHARMA provides simulation-based standardized residual diagnostics specifically designed for generalized linear mixed models and hierarchical regression models. Unlike conventional residual diagnostics that rely primarily on visual inspection of residual plots, histograms, and Q-Q plots, DHARMA generates randomized quantile residuals that are less influenced by model complexity and non-normal data distributions. This makes it particularly suitable and more reliable for evaluating mixed-effects models commonly used in auditory and behavioral research. Residual uniformity was assessed using

the Kolmogorov–Smirnov test (Massey, 1951), while nonparametric dispersion tests were performed to evaluate over-dispersion and under-dispersion. Homogeneity of variance was examined using Levene’s test (Levene, 1960).

Secondary Task

Secondary task performance was analyzed using a linear mixed-effects regression model (LMER) implemented in RStudio using the lme4 test package (Bates et al., 2015). A LMER was employed because the dependent variable consisted of continuous listening effort scores obtained repeatedly across RTs and SNRs conditions, across groups. Linear mixed-effects models are particularly advantageous for repeated-measures auditory experiments because they account for within-subject dependencies while simultaneously evaluating the influence of fixed experimental factors and their interactions. Such models have been extensively used in listening effort and speech-in-noise research to evaluate the combined influence of acoustic degradation, linguistic complexity, and speaker characteristics on behavioral outcomes (Downey et al., 2026; Johns et al., 2024; Rovetti et al., 2023).

The model was estimated using Restricted Maximum Likelihood (REML), which provides robust estimation of variance components in repeated-measures designs. The model was employed to examine the fixed effects of Reverberation Time ($RT_{min} = 0.74$ s and $RT_{max} = 4.37$ s), Signal-to-Noise Ratio (SNR: 0 dB and 4 dB), and Group (Hindi, Malayalam, and Kannada) along with their interactions on listening effort performance measured through the secondary recall task. Participant ID was included as a random intercept to account for inter-individual variability across repeated observations.

Model diagnostics were carried out using the DHARMa (Diagnostics for Hierarchical Regression Models using Randomized Quantile Residuals) package in R.

Residual uniformity was assessed using the Kolmogorov–Smirnov test, while dispersion diagnostics were performed to evaluate over-dispersion and under-dispersion. Outlier analysis was conducted using a binomial outlier test. Visual inspection of residual plots, including Q-Q plots and histograms, was also performed to assess normality and homoscedasticity assumptions. In addition, the correlation matrix of fixed effects was examined to assess potential multicollinearity among predictors.

NASA-TLX

Subjective listening effort using NASA-TLX was analyzed using a Linear Mixed-Effects Regression (LMER) model implemented in RStudio using the lme4 package (Bates et al., 2015). The model was employed to examine the effects of Reverberation Time (RT minimum = 0.74 s, RT average = 1.66 s, and RT maximum = 4.37 s), Signal-to-Noise Ratio (SNR: 0 dB and 4 dB), and Group (Hindi, Malayalam, and Kannada) on subjective listening effort measured using the NASA Task Load Index (NASA-TLX). Participant ID was included as a random intercept to account for within-subject variability across repeated observations.

The model was estimated using Restricted Maximum Likelihood (REML), which provides robust estimation of variance components in repeated-measures designs. The fixed effects structure included RT, SNR, and Group along with their interaction terms to examine both the independent and combined effects of acoustic and linguistic variables on subjective listening effort.

Model diagnostics were carried out using graphical inspection methods and the DHARMa (Diagnostics for Hierarchical Regression Models using Randomized Quantile

Residuals) package in R. Residual plots, histograms, and Q-Q plots were visually inspected to assess the assumptions of normality and homogeneity of variance. Residual uniformity was assessed using the Kolmogorov–Smirnov test, while dispersion diagnostics were performed to evaluate over-dispersion and under-dispersion. Outlier analysis was conducted using a binomial outlier test to determine whether any observations exerted undue influence on the model.

Correlational analysis

The relationship between language proficiency and listening effort measures was examined using Spearman’s rank-order correlation analysis in R Studio. Correlations were computed between LEAP-Q scores and (1) primary task scores, (2) secondary task scores, and (3) NASA-TLX scores across different RT and SNR conditions. Since the data did not consistently meet assumptions of normality and linearity, the non-parametric Spearman correlation coefficient (S_p) was used.

The dataset was imported into R Studio and variables were cleaned and recoded prior to analysis. Reverberation conditions were categorized as Minimum RT, Average RT, and Maximum RT, while noise conditions were categorized as 0 dB SNR and 4 dB SNR. Primary and secondary task scores were converted into long format to facilitate grouped correlation analysis across task conditions.

Spearman correlation coefficients and corresponding p-values were calculated separately for each combination of RT, SNR, and task condition. Statistical significance was interpreted at $p < .05$, $p < .01$, and $p < .001$ levels.

The obtained correlation coefficients were further visualized using heat maps generated with the ggplot2 package in R. Heat maps displayed the direction and strength of correlations across RT and SNR conditions using a diverging colour gradient, where positive

and negative correlations were represented by contrasting colours and non-available conditions were indicated separately. Correlation tables and heat maps were exported and saved for further interpretation and reporting.

Test-retest analysis

Test-retest reliability of the experimental measures was assessed using Cronbach's alpha analysis in R Studio. Reliability analysis was performed separately for primary task scores, secondary task scores, and NASA-TLX scores obtained during the test and retest sessions across different RT, SNR and Group conditions.

Prior to analysis, the dataset was cleaned and organized in R Studio, and scores from test and retest sessions were paired for each participant across all listening conditions. Cronbach's alpha coefficients (α) were calculated to evaluate the internal consistency and stability of the measures across repeated administrations. Higher alpha values indicated greater consistency between test and retest performances.

Reliability coefficients were interpreted based on commonly accepted criteria, where α values above 0.70 indicated acceptable reliability, values above 0.80 indicated good reliability, and values above 0.90 indicated excellent reliability. Descriptive statistics including means and standard deviations for test and retest scores were also computed to support interpretation of the reliability findings. The obtained reliability coefficients were summarized across all listening conditions to examine the consistency of behavioral and subjective listening effort measures under varying RT, SNR and Group conditions.

CHAPTER 4

RESULTS

The study aimed to investigate the effects of native language-based experience on understanding speech in noise. Listening effort using the dual-task paradigm (primary and secondary tasks) was examined as a function of acoustic conditions (2 signal to noise ratio - SNRs and 3 reverberation times- RTs) and language characteristics of the speakers (Groups). The behavioural task scores of speech perception i.e. primary task scores; of cognition i.e. secondary task scores, and subjective ratings (NASA Task Load Index) scores were analysed. The results will be discussed under the following headings.

4.1 Overall effects of Groups, RT, and SNR on Primary, Secondary Tasks.

4.2. Effects of RT, SNR, and Groups on Primary Task

4.2.1 Effect of Language Group on Primary Task Performance

4.1.2 Effect of Reverberation Time (RT) on Primary Task Performance

4.1.3 Effect of Signal-to-Noise Ratio (SNR) on Primary Task Performance

4.3. Effects of RT, SNR, and Groups on Secondary Task

4.3.1 Effect of Language Group on Secondary Task Performance

4.3.2 Effect of Reverberation Time (RT) on Secondary Task Performance

4.3.3 Effect of Signal-to-Noise Ratio (SNR) on Secondary Task Performance

4.4 Effect of RT, SNR and Group on Subjective Listening Effort (NASA-TLX)

4.4.1 Effect of Language Group on NASA-TLX Scores

4.4.2 Effect of Reverberation Time (RT) on NASA-TLX Scores

4.4.3 Effect of Signal-to-Noise Ratio (SNR) on NASA-TLX Scores

4.5 Correlation between LEAP-Q Scores and Listening Effort Measures

4.5.1. Correlation between LEAP-Q Scores with Primary and Secondary Task Performance

4.5.2 Correlation between LEAP-Q Scores and NASA-TLX Ratings

4.6 Results of test-retest reliability

4.1 Overall effects of Groups, RT, and SNR on Primary, Secondary Tasks.

Median and interquartile range (errors bars) scores across SNRs, RTs and language groups revealed several consistent patterns as shown in the Figure 4.1. Across all reverberation conditions, task scores were higher at SNR 4 compared to SNR 0, indicating improved performance under more favourable listening conditions.

A clear effect of reverberation time was observed, with the highest scores obtained under minimum reverberation, followed by average, and the lowest scores under maximum reverberation. This trend was evident for both primary and secondary tasks, suggesting that increased reverberation negatively impacts speech perception performance.

Group-wise comparisons indicated that the Kannada group consistently demonstrated the highest task scores, followed by the Malayalam group, while the Hindi group showed the lowest performance across conditions. These differences were more pronounced under adverse listening conditions, particularly at lower SNR and higher reverberation levels.

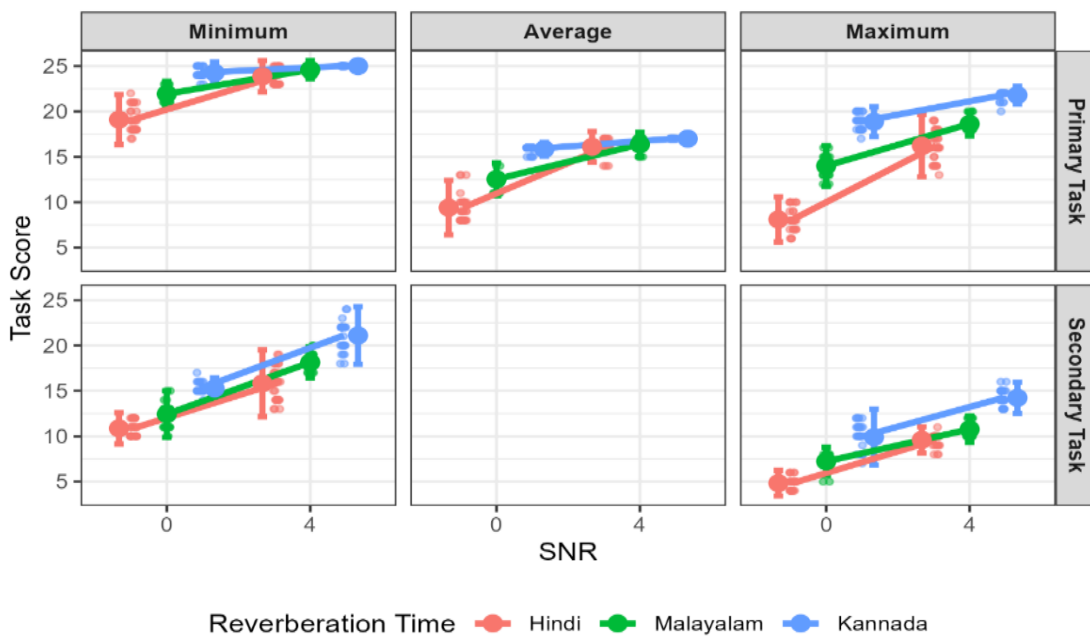
Furthermore, primary task scores were consistently higher than secondary task scores, with the latter showing greater sensitivity to changes in SNR and reverberation. This

pattern suggests increased cognitive load and listening effort under more challenging acoustic environments.

Overall, the descriptive findings indicate that both acoustic factors (SNR and reverberation) and speaker characteristics (language background) jointly influence task performance, with non-native speakers being more adversely affected under degraded listening conditions.

Figure-4.1

Interaction plots illustrating primary and secondary task performance across Signal-to-Noise Ratio (SNR) conditions, Reverberation Time (Minimum, Average, and Maximum RT), and speaker groups (Hindi, Malayalam, and Kannada). Coloured lines represent the median task scores across conditions. Individual dots indicate each participant performance and error bars represent the interquartile range (IQR).



4.2 Effects of RT, SNR, and Groups on Primary Task

The effects of RT, SNR, and language group were examined using beta generalized linear mixed model (GLMM) in Rstudio. Prior to examining the effects of RT, SNR, and language group on primary task performance, diagnostic analyses were conducted to evaluate the adequacy of the beta GLMM.

Residual diagnostics using the Diagnostics for Hierarchical Regression Models using Randomized Quantile Residuals (DHARMA) indicated that the assumptions of the beta GLMM were reasonably satisfied. Residual uniformity assessed using the Kolmogorov–Smirnov test revealed a small but statistically significant deviation from the expected distribution ($D = 0.07, p = .01$). However, the deviation was minimal, and visual inspection of residual distributions suggested that the model adequately fit the data. The nonparametric dispersion test demonstrated no evidence of over-dispersion or under-dispersion (dispersion = 1.03, $p = .69$), indicating that the variance structure of the model was appropriate for the present data. Residual diagnostics across SNR conditions also revealed only minor deviations in residual uniformity [0 dB: $D = 0.09, p = .02$; SNR2: $D = 0.10, p = .007$]. Levene’s test for homogeneity of variance was non-significant [$F(1, 538) = 2.61, p = .11$], suggesting that the assumption of homogeneity of variance was adequately met. Furthermore, visual inspection of Q-Q plots and residual histograms indicated no substantial departures from model assumptions. Overall, these diagnostic analyses suggested that the beta mixed-effects generalized linear model provided an acceptable fit to the primary task data and was appropriate for examining the main and interaction effects of RT, SNR, and language group on speech recognition performance. In addition, the random-effects analysis demonstrated variability across participants, supporting the inclusion of participant-specific random intercepts in the model. The standard deviation of the random

intercept for participants was 0.08 (95% CI [0.05, 0.11]), indicating moderate between-subject variability in speech recognition performance.

The beta GLMM model was built using RT, SNR, and language group as fixed effect, and Participant_ID as random effect. The model is represented using equation 1:

$$\text{Primary task scores} \sim \text{RT} * \text{SNR} * \text{Group} + (1 | \text{Participant_ID}) - \text{eq (1)}$$

The beta GLMM revealed significant fixed effects of RT, SNR, and language group on primary task performance as shown in Table 4.1. The fixed-effects analysis revealed significant effects of RT, SNR, and language group, and their interactions on primary task performance. Effect sizes were interpreted based on the magnitude of the standardized beta coefficients.

Table 4.1

Fixed-effect estimates from the Beta GLMM showing the effects of group, RT, SNR, and on primary task scores.

Parameter	Coefficient (β)	Std. Error	95% CI [Lower bound, Upper bound]	<i>z</i> value	<i>p</i> value
(Intercept)	1.17	0.04	[1.08, 1.25]	27.11	< .001
Avg RT	-1.67	0.05	[-1.77, -1.56]	-30.72	< .001
Max RT	-1.89	0.06	[-2.00, -1.79]	-34.37	< .001
4 dB SNR	2.09	0.10	[1.90, 2.28]	21.62	< .001
Malayalam	0.76	0.07	[0.62, 0.89]	11.02	< .001
Kannada	2.37	0.11	[2.16, 2.58]	21.98	< .001

Avg RT × 4 dB SNR	-1.00	0.11	[-1.22, -0.79]	-9.18	< .001
Max RT × 4 dB SNR	-0.75	0.11	[-0.96, -0.53]	-6.82	< .001
Avg RT × Malayalam	-0.25	0.08	[-0.41, -0.09]	-3.01	0.003
Max RT × Malayalam	0.21	0.08	[0.05, 0.37]	2.52	0.012
Avg RT × Kannada	-1.33	0.12	[-1.56, -1.10]	-11.32	< .001
Max RT × Kannada	-0.54	0.12	[-0.77, -0.30]	-4.51	< .001
4 dB SNR × Malayalam	0.03	0.16	[-0.28, 0.35]	0.19	0.84
4 dB SNR × Kannada	-0.98	0.20	[-1.37, -0.59]	-4.95	< .001
(Avg RT × 4 dB SNR) × Malayalam	-0.49	0.18	[-0.84, -0.15]	-2.80	0.005
(Max RT × 4 dB SNR) × Malayalam	-0.56	0.18	[-0.90, -0.21]	-3.13	0.002
(Avg RT × 4 dB SNR) × Kannada	0.09	0.21	[-0.32, 0.51]	0.45	0.65
(Max RT × 4 dB SNR) × Kannada	0.40	0.21	[-0.02, 0.82]	1.85	0.06

Note: SNR -signal to noise ratio ; RT-reverberation time ; Min -minimum ; Avg-average ; Max –maximum

Visual inspection of Table 4.1 revealed a very large negative effect of RT, indicating that speech recognition performance significantly decreased as reverberation increased. Compared to the reference condition (Minimum RT), both Average RT ($\beta = -1.67$, 95% CI [-1.77, -1.56], $z = -30.72$, $p < .001$) and Maximum RT ($\beta = -1.89$, 95% CI [-2.00, -1.79], $z = -34.37$, $p < .001$) demonstrated very large detrimental effects on primary task scores. Examination of effect size as indicated by ' β ' revealed very large negative effects for both

Average RT and Maximum RT conditions, implying that increased reverberation substantially reduced speech recognition performance and increased listening difficulty across participants.

SNR also showed a very large positive effect, with significantly better performance at 4 dB SNR compared to 0 dB SNR ($\beta = 2.09$, 95% CI [1.90, 2.28], $z = 21.62$, $p < .001$). Examination of effect size revealed a very large positive effect, implying that improved acoustic clarity through higher SNR considerably enhanced speech recognition performance and reduced the adverse effects of background noise.

Language of the speaker group had a substantial influence on performance. Relative to the Hindi group, the Malayalam group showed a moderate positive effect ($\beta = 0.76$, 95% CI [0.62, 0.89], $z = 11.02$, $p < .001$), whereas the Kannada group demonstrated a very large positive effect ($\beta = 2.37$, 95% CI [2.16, 2.58], $z = 21.98$, $p < .001$), indicating superior speech recognition performance among native Kannada speakers. Examination of effect size revealed a moderate effect for the Malayalam group and a very large effect for the Kannada group, implying that familiarity with the target language strongly facilitated speech perception, particularly for native Kannada speakers.

Comparison of the effect sizes revealed that reverberation time (RT) exerted the strongest influence on primary task performance, followed by signal-to-noise ratio (SNR), and then speaker group. RT demonstrated the largest effect sizes overall, particularly under Maximum RT conditions, indicating that increased reverberation had the greatest detrimental impact on speech recognition performance. SNR showed the second largest effect, suggesting that improvements in acoustic clarity substantially enhanced performance, although its influence was comparatively smaller than reverberation. Speaker group demonstrated relatively smaller but still substantial effects, with native Kannada speakers consistently outperforming non-native speakers

The interaction between RT and SNR demonstrated large negative effects for both Average RT $\times$ 4 dB SNR ($\beta = -1.00$, 95% CI $[-1.22, -0.79]$, $z = -9.18$, $p < .001$) and Maximum RT $\times$ 4 dB SNR ($\beta = -0.75$, 95% CI $[-0.96, -0.53]$, $z = -6.82$, $p < .001$), suggesting that the influence of reverberation varied across SNR conditions. Interactions between RT and speaker group also revealed significant effects. Average RT $\times$ Malayalam showed a small negative effect ($\beta = -0.25$, $p = .003$), whereas Maximum RT $\times$ Malayalam showed a small positive effect ($\beta = 0.21$, $p = .012$). In contrast, Average RT $\times$ Kannada demonstrated a very large negative effect ($\beta = -1.33$, $p < .001$), while Maximum RT $\times$ Kannada showed a moderate negative effect ($\beta = -0.54$, $p < .001$). The interaction between 4 dB SNR and Malayalam was negligible and non-significant ($\beta = 0.03$, $p = .846$), whereas 4 dB SNR $\times$ Kannada demonstrated a large negative effect ($\beta = -0.98$, $p < .001$). Examination of effect sizes revealed that reverberation and SNR interacted differently across speaker groups, implying that adverse acoustic conditions disproportionately affected non-native speakers, while native Kannada speakers demonstrated greater resilience under challenging listening conditions.

Among the three-way interactions, Average RT $\times$ 4 dB SNR $\times$ Malayalam ($\beta = -0.49$, $p = .005$) and Maximum RT $\times$ 4 dB SNR $\times$ Malayalam ($\beta = -0.56$, $p = .002$) showed moderate negative effects, implying that Malayalam speakers experienced additional reductions in performance under combined reverberation and noise conditions. However, the interactions involving Kannada speakers were non-significant, with Average RT $\times$ 4 dB SNR $\times$ Kannada showing a negligible effect ($\beta = 0.09$, $p = .656$) and Maximum RT $\times$ 4 dB SNR $\times$ Kannada demonstrating a small positive effect ($\beta = 0.40$, $p = .064$). These findings imply that native Kannada speakers were comparatively less affected by the combined acoustic degradations of reverberation and background noise.

The main effects of RT, SNR, and speaker group on primary task performance were also confirmed using Type II Wald χ^2 tests obtained from the beta regression GLMM. The analysis revealed highly significant main effects of RT, [$\chi^2(2) = 3171.08, p < .001$], SNR, [$\chi^2(1) = 2282.39, p < .001$], and speaker group, [$\chi^2(2) = 1209.71, p < .001$], as shown in Table 4.2.

Table 4.2

ANOVA Results (Type II Wald χ^2 Tests) for Primary Task Performance. Corresponding effect sizes are also mentioned.

Effect	χ^2	df	p-value	Effect Size (η_p^2)
RT	3171.08	2	< .001	-1.09×10^{-7}
SNR	2282.39	1	< .001	7.15×10^{-8}
Group	1209.71	2	< .001	-1.10×10^{-8}
RT $\times$ SNR	142.99	2	< .001	4.27×10^{-8}
RT $\times$ Group	492.69	4	< .001	1.45×10^{-8}
SNR $\times$ Group	389.58	2	< .001	3.70×10^{-8}
RT $\times$ SNR $\times$ Group	27.24	4	< .001	-9.35×10^{-5}

Note: SNR -signal to noise ratio; RT-reverberation time

As shown in Table 4.2, a significant negative effect of RT was seen, indicating that increasing reverberation adversely affected primary task performance. Speech recognition scores declined as reverberation increased from Minimum reverberation to Maximum reverberation, suggesting that reverberant listening conditions substantially degrade speech perception and increase perceptual demands.

A similarly strong main effect of SNR was observed, demonstrating that participants performed significantly better at the higher SNR condition (+4 dB) compared to the poorer SNR condition (0 dB). This finding indicates that improved signal clarity facilitated more accurate recognition of the target words.

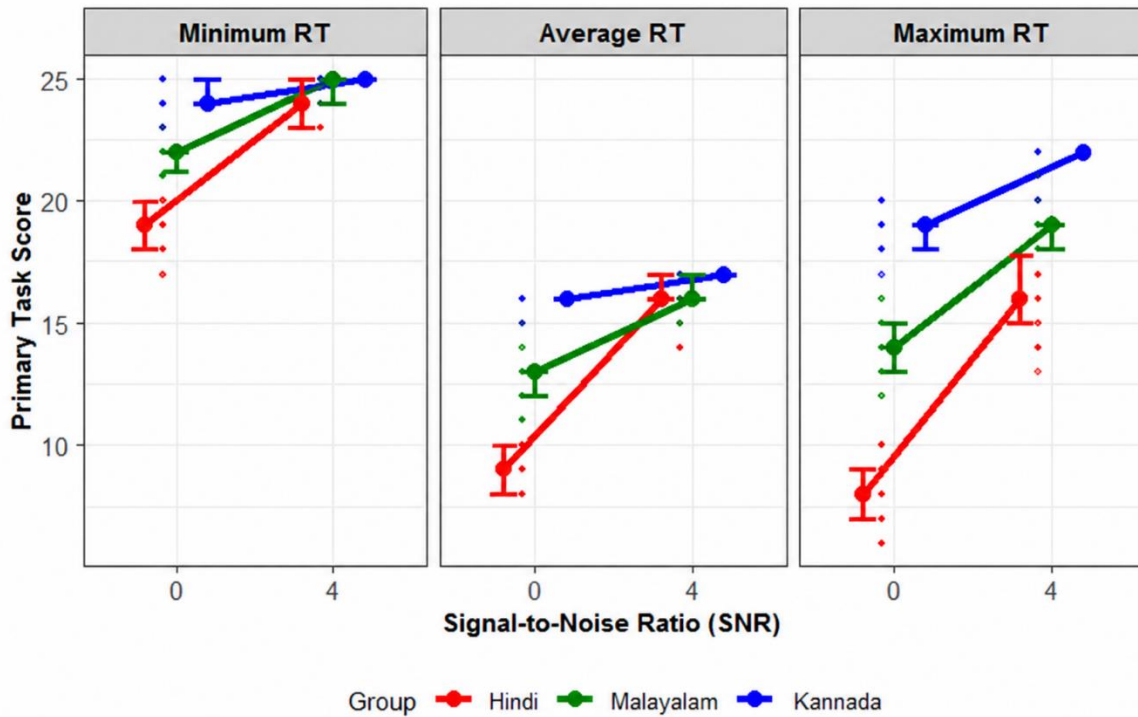
Further examination of interaction effects revealed a statistically significant RT $\times$ SNR interaction, [$\chi^2(2) = 142.99, p < .001, \eta_p^2 = 4.27 \times 10^{-8}$], indicating that the effect of reverberation varied across SNR conditions. The RT $\times$ Group interaction was also significant, [$\chi^2(4) = 492.69, p < .001, \eta_p^2 = 1.45 \times 10^{-8}$], suggesting that the influence of reverberation differed across speaker groups. Similarly, the SNR $\times$ Group interaction was statistically significant, [$\chi^2(2) = 389.58, p < .001, \eta_p^2 = 3.70 \times 10^{-8}$], indicating that the benefit obtained from improved SNR varied among the groups. Importantly, the three-way interaction between RT, SNR, and Group was statistically significant, [$\chi^2(4) = 27.24, p < .001, \eta_p^2 = 9.35 \times 10^{-5}$]. This finding suggests that the combined effects of reverberation and noise on speech recognition differed across native and non-native speaker groups.

4.2.1 Effect of Language Group on Primary Task Performance

As Beta GLMM model revealed significant main effect of Groups, as well as its significant interaction with SNR and RT (Table 4.1), pairwise comparison with Bonferroni correction were carried out for comparison of primary task scores of three Groups across two SNRs and three RT conditions. Figure 4.2 shows interaction plots with median and IQR along with individual data points for primary task as a function of Group.

Figure 4.2

Interaction plots illustrating primary task performance across SNR conditions under Minimum, Average, and Maximum RT conditions for Hindi, Malayalam, and Kannada speaker groups. Coloured lines represent the median primary task scores for each speaker group across SNR conditions. Individual dots indicate each participant performance and Error bars represent the interquartile range (IQR).



Inspection of Figure 4.2 reveals clear Group-wise simple effects at each SNR $\times$ RT combination. Under Minimum RT at 0 dB SNR, Kannada speakers demonstrated the highest primary task scores, followed by Malayalam speakers, while Hindi speakers showed the lowest performance. At 4 dB SNR under Minimum RT, all groups showed improved performance; however, Kannada speakers continued to demonstrate the highest scores, with Hindi and Malayalam speakers showing comparable performance levels.

Under Average RT at 0 dB SNR, Kannada speakers again achieved the highest primary task scores, followed by Malayalam speakers, whereas Hindi speakers exhibited the lowest scores. At 4 dB SNR under Average RT, performance improved across all groups,

and the difference between Hindi and Malayalam speakers became smaller, although Kannada speakers continued to outperform both groups.

Under Maximum RT at 0 dB SNR, the group differences were more pronounced, with Kannada speakers maintaining the highest performance, Malayalam speakers demonstrating intermediate scores, and Hindi speakers showing the lowest scores. At 4 dB SNR under Maximum RT, all groups demonstrated improved performance relative to 0 dB SNR; however, Kannada speakers continued to achieve the highest scores, followed by Malayalam speakers and Hindi speakers.

Across all listening conditions, native Kannada speakers (Group 3) consistently demonstrated the highest speech recognition performance, followed by Malayalam non-native speakers (Group 2), while Hindi non-native speakers (Group 1) exhibited the lowest performance. These findings indicate that linguistic familiarity and language proficiency significantly influenced speech recognition under adverse acoustic conditions.

Post hoc pairwise comparisons performed using estimated marginal means (EMMs) with Bonferroni adjustment for multiple comparisons, showed significant differences in primary task scores between Group conditions as shown in Table 4.3. Comparisons were performed on the log odds ratio scale, and results were reported as z-values and Bonferroni-adjusted p-values. Larger absolute z-values indicated stronger differences between conditions, with $p < .05$ considered statistically significant. Pairwise comparisons revealed significant differences across groups for most RT and SNR conditions. The largest differences were consistently observed between Group 1 (Hindi) and Group 3 (Kannada), demonstrating a strong native-language advantage during speech recognition tasks.

Table 4.3

Pairwise Group Comparisons on Primary Task Performance Across RT and SNR Conditions

SNR	RT	Hindi vs Malayalam			Hindi vs Kannada			Malayalam vs Kannada		
		Odd's ratio	SE	z, p	Odd's ratio	SE	z, p	Odd's ratio	SE	z, p
0 dB	Min RT	5.28	0.31	27.91, < .001	6.61	0.40	30.90, < .001	1.25	0.06	5.05, < .001
	Avg RT	0.34	0.01	30.65, < .001	0.27	0.01	34.56, < .001	0.81	0.03	5.73, < .001
	Max RT	0.02	0.00	30.46, < .001	0.67	0.02	22.39, < .001	0.04	0.01	17.24, < .001
4 dB	Min RT	12.75	1.58	20.56, < .001	12.39	1.57	19.85, < .001	0.97	0.04	-0.74, 1.000
	Avg RT	0.34	0.01	30.65, < .001	0.33	0.01	24.94, < .001	0.66	0.03	22.37, < .001
	Max RT	0.26	0.01	28.69, < .001	0.44	0.02	17.08, < .001	0.47	0.03	12.99, < .001

Note: SNR -signal to noise ratio ; RT-reverberation time ; Min -minimum ; Avg-average ; Max -maximum

Interaction Effects Involving Group

The interaction between RT and Group was statistically significant, [$\chi^2(4) = 492.69$, $p < .001$, $\eta_p^2 = 1.45 \times 10^{-8}$], indicating that the effect of reverberation differed across speaker groups. Examination of effect size revealed a measurable interaction effect, suggesting that non-native speakers demonstrated greater reductions in performance with increasing reverberation compared to native Kannada speakers.

Similarly, the SNR $\times$ Group interaction was statistically significant, [$\chi^2(2) = 389.58$, $p < .001$, $\eta_p^2 = 3.70 \times 10^{-8}$], suggesting that the speaker groups differed in the extent to which they benefited from improved SNR conditions. The observed effect size indicated that linguistic background influenced the degree to which speakers benefited from improved acoustic clarity.

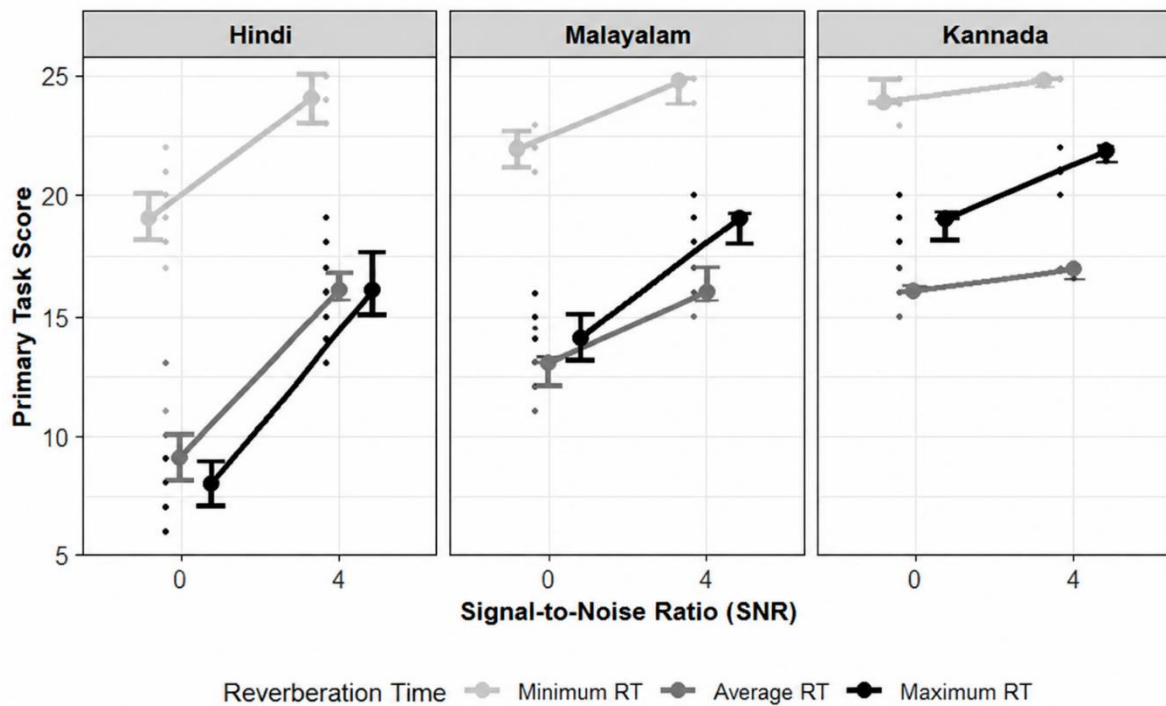
Importantly, the three-way interaction between RT, SNR, and Group was statistically significant, [$\chi^2(4) = 27.24$, $p < .001$, $\eta_p^2 = -9.35 \times 10^{-5}$]. This finding indicates that the combined effects of reverberation and noise varied as a function of linguistic background, with non-native speakers experiencing greater speech perception difficulty under acoustically adverse conditions. The comparatively larger effect size observed for the three-way interaction further suggests that the combined influence of reverberation and noise differed substantially across speaker groups.

4.1.2 Effect of Reverberation Time (RT) on Primary Task Performance

As Beta GLMM revealed significant main effect of RT as well as its significant interaction with Group and SNR (Table 4.1), pairwise comparison with Bonferroni correction were carried out for comparison of Primary Task scores of three RT conditions for each group across SNRs. Figure 4.3 shows interaction plots with median and IQR along with individual data points for primary task as a function of RT.

Figure 4.3

Interaction plots illustrating primary task performance across SNR conditions for the Hindi, Malayalam, and Kannada speaker groups under Minimum, Average, and Maximum RT conditions. Coloured lines represent the median primary task scores across SNR conditions for each reverberation condition. Individual dots indicate each participant performance and Error bars represent the interquartile range (IQR).



Inspection of Figure 4.3 revealed clear RT-wise simple effects across each SNR $\times$ Group combination. At 0 dB SNR, the Hindi group demonstrated the highest primary task scores under Minimum RT, followed by Average RT, whereas the lowest scores were observed under Maximum RT, reflecting a decline in speech recognition performance with increasing reverberation. A similar RT-related trend was observed in the Malayalam group, where performance progressively decreased from Minimum RT to Maximum RT. The

Kannada group also exhibited the same pattern at 0 dB SNR; however, their overall scores remained consistently higher than those of the non-native speaker groups.

At 4 dB SNR, the Hindi group again demonstrated superior performance under Minimum RT, intermediate performance under Average RT, and the lowest performance under Maximum RT. The Malayalam group similarly showed a progressive decline in primary task scores with increasing reverberation. Likewise, the Kannada group at 4 dB SNR exhibited the highest scores under Minimum RT and the lowest scores under Maximum RT, while continuing to demonstrate overall better speech recognition performance compared to the Hindi and Malayalam groups.

Pairwise comparisons demonstrated significant differences between Minimum RT and Average RT, Minimum RT and Maximum RT, as well as Average RT and Maximum RT across most listening conditions, as shown in Table 4.4. For Hindi group, the comparison between Minimum RT at 0 dB SNR and Average RT at 0 dB SNR was statistically significant ($z = 27.91, p < .001$), while Minimum RT at 0 dB SNR also differed significantly from Maximum RT at 0 dB SNR ($z = 30.90, p < .001$). Similar RT-related trends were observed for Malayalam group and Kannada group across both SNR conditions.

Table 4.4

Pairwise Comparison of RT Conditions on Primary Task Performance

SNR	RT	Hindi			Malayalam			Kannada		
Condition										
		Odds	SE	z, p	Odds	SE	z, p	Odds	SE	z, p
		Ratio			Ratio			Ratio		

0 dB	Min	vs	5.28	0.31	27.91,	6.67	0.47	26.71,	18.56	2.28	23.80,
	Avg				< .001			< .001			< .001
	Min	vs	6.61	0.40	30.90,	5.28	0.38	23.16,	10.51	1.31	18.94,
	Max				< .001			< .001			< .001
	Avg	vs	1.25	0.06	5.05,	0.79	0.03	-5.46,	0.57	0.03	-12.12,
	Max				< .001			< .001			< .001
4 dB	Min	vs	12.75	1.58	20.56,	23.68	3.72	20.15,	34.03	6.16	19.49,
	Avg				< .001			< .001			< .001
	Min	vs	12.39	1.57	19.85,	15.52	2.48	17.14,	11.06	2.06	12.94,
	Max				< .001			< .001			< .001
	Avg	vs	0.97	0.04	-0.74,	0.66	0.03	-10.33,	0.33	0.02	-22.37,
	Max				1.000			< .001			< .001

Note: SNR -signal to noise ratio ; RT-reverberation time ; Min -minimum ; Avg-average ; Max -maximum

Overall, across all speaker groups and SNR conditions, primary task performance was consistently highest under Minimum RT, followed by Average RT, and lowest under Maximum RT. These findings indicate that increasing reverberation progressively degraded speech recognition performance irrespective of speaker group or SNR condition.

Interaction Effects Involving RT

The interaction between RT and SNR was statistically significant [$\chi^2(2) = 142.99, p < .001, \eta_p^2 = 4.27 \times 10^{-8}$], indicating that the effect of reverberation differed across SNR conditions. Examination of effect size revealed a measurable interaction effect, suggesting that the detrimental effect of reverberation became more pronounced under poorer SNR conditions.

The $RT \times Group$ interaction was also statistically significant [$\chi^2(4) = 492.69, p < .001, \eta_p^2 = 1.45 \times 10^{-8}$], suggesting that reverberation affected the speaker groups differently. Although all groups showed reduced performance with increasing RT, the non-native groups demonstrated greater reductions in performance under Average RT and Maximum RT conditions. The observed effect size further indicated that reverberation had a differential impact across speaker groups.

Further, the three-way interaction between RT, SNR, and Group was statistically significant, [$\chi^2(4) = 27.24, p < .001, \eta_p^2 = 9.35 \times 10^{-5}$], indicating that the combined effects of reverberation and noise differed across speaker groups. Non-native speakers demonstrated greater speech recognition difficulty under acoustically adverse conditions. The comparatively larger effect size observed for the three-way interaction suggests that linguistic background substantially influenced the combined effects of reverberation and noise on speech recognition performance.

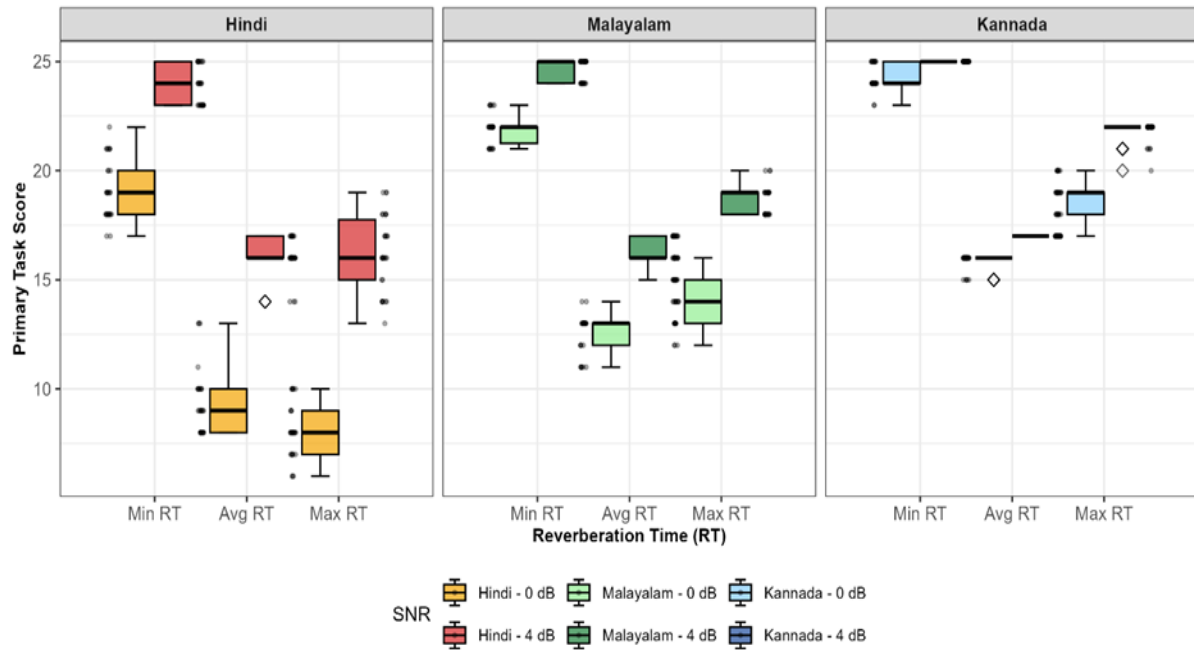
4.1.3 Effect of Signal-to-Noise Ratio (SNR) on Primary Task Performance

As Beta GLMM revealed significant main effect of SNR, as well as its significant interaction with Group and RT (Table 4.1), pairwise comparison with Bonferroni correction were carried out for comparison of Primary Task scores of two SNR conditions for each group across RTs. Figure 4.4 shows boxplots with median and IQR along with individual data points for primary task as a function of SNR.

Figure 4.4

Boxplots illustrating primary task performance across RT conditions (Minimum, Average, and Maximum RT) for Hindi, Malayalam, and Kannada speaker groups under 0 dB and 4

dB SNR conditions. The central horizontal line within each box represents the median score, while the box boundaries indicate the interquartile range (IQR). Whiskers represent the spread of the data beyond the quartiles, and individual dots indicate each participant-performance and outliers.



Inspection of Figure 4.4 revealed clear SNR-wise simple effects at each RT $\times$ Group combination. Under Minimum RT for the Hindi group, primary task scores were higher at 4 dB SNR compared to 0 dB SNR, indicating improved speech recognition performance under more favourable listening conditions. A similar pattern was observed for the Malayalam group under Minimum RT, where scores at 4 dB SNR exceeded those at 0 dB SNR, reflecting better speech perception with improved SNR. For the Kannada group under Minimum RT, the increase from 0 dB to 4 dB SNR was also evident, with the highest overall scores observed at 4 dB SNR among all groups.

Under Average RT for the Hindi group, primary task scores again showed an increase at 4 dB SNR relative to 0 dB SNR, although the overall scores were lower compared to

Minimum RT. The Malayalam group under Average RT demonstrated a similar improvement with increasing SNR, with higher scores at 4 dB SNR than at 0 dB SNR. Likewise, for the Kannada group under Average RT, primary task performance improved at 4 dB SNR, maintaining the same directional trend despite the negative influence of reverberation.

Under Maximum RT for the Hindi group, primary task scores remained higher at 4 dB SNR compared to 0 dB SNR, though performance was substantially reduced relative to the other RT conditions. The Malayalam group under Maximum RT also demonstrated improved scores at 4 dB SNR relative to 0 dB SNR. Similarly, the Kannada group under Maximum RT showed higher primary task scores at 4 dB SNR, while still maintaining overall superior performance compared to the other speaker groups.

Across all RT and group conditions, primary task performance was consistently higher at 4 dB SNR compared to 0 dB SNR. Improved signal clarity facilitated better identification of the target words during the primary task.

Pairwise comparisons further demonstrated significant differences between 0 dB and 4 dB SNR conditions across all speaker groups as shown in Table 4.5. For Group 1(Hindi), minimum RT SNR 0 dB differed significantly from minimum RT SNR 4 dB ($z = -14.86, p < .001$), while similar significant improvements were observed for Group 2 (Malayalam) ($z = -11.23, p < .001$) and Group 3 (Kannada) ($z = -3.77, p = .0025$).

These findings demonstrate that increasing SNR significantly improved primary task performance across all listening conditions.

Table 4.5

Pairwise Comparison of SNR Conditions (0 dB vs 4 dB) on Primary Task Performance

SNR	RT	Hindi			Malayalam			Kannada		
		Odds	SE	z, p	Odds	SE	z, p	Odds	SE	z, p
		Ratio			Ratio			Ratio		
0 dB	Min	0.14	0.02	-14.86,	0.15	0.03	-11.23,	0.44	0.10	-3.77,
vs 4				< .001			< .001			.0025
0 dB	Avg	0.34	0.01	-30.65,	0.53	0.02	-17.99,	0.81	0.03	-5.73, <
vs 4				< .001			< .001			.001
0 dB	Max	0.26	0.01	-28.69,	0.44	0.02	-17.08,	0.47	0.03	-12.99,
vs 4				< .001			< .001			< .001

Note: SNR -signal to noise ratio ; RT-reverberation time ; Min -minimum ; Avg-average ; Max -maximum

Interaction Effects Involving SNR

The $RT \times SNR$ interaction was statistically significant, [$\chi^2(2) = 142.99, p < .001, \eta_p^2 = 4.27 \times 10^{-8}$], indicating that the benefit obtained from improved SNR varied across reverberation conditions. Specifically, improvements in SNR resulted in greater performance gains under lower reverberation conditions compared to highly reverberant conditions. Examination of effect size revealed a measurable interaction effect, suggesting that reverberation substantially moderated the influence of SNR on speech recognition performance.

The $SNR \times Group$ interaction was also statistically significant, [$\chi^2(2) = 389.58, p < .001, \eta_p^2 = 3.70 \times 10^{-8}$], suggesting that the speaker groups differed in the extent to which

they benefited from improved signal clarity. Native Kannada speakers consistently demonstrated superior performance across SNR conditions. Examination of effect size indicated that linguistic background meaningfully influenced the extent of benefit obtained from improved acoustic clarity.

Additionally, the three-way interaction between RT, SNR, and Group was statistically significant, [$\chi^2(4) = 27.24, p < .001 \eta_p^2 = -9.35 \times 10^{-5}$], indicating that the combined effects of reverberation and noise differed across speaker groups. Examination of effect size further revealed that the joint influence of reverberation and SNR varied considerably as a function of linguistic background, with non-native speakers demonstrating greater speech perception difficulty under acoustically adverse conditions.

4.3 Effect of RT, SNR and Groups on Secondary Task

The effects of RT, SNR, and language group were examined using a linear mixed-effects regression model (LMER) implemented in RStudio using the lme4 and lmerTest packages (Bates et al., 2015). Prior to examining the effects of RT, SNR, and language group on secondary task performance, diagnostic analyses were conducted to evaluate the adequacy of LMER model.

Residual diagnostics using the DHARMA indicated that the assumptions of the LMER model were reasonably satisfied. Residual uniformity assessed using the Kolmogorov–Smirnov test indicated no significant deviation from the expected distribution ($D = 0.093, p = .093$), suggesting that the residuals were approximately normally distributed. Scaled residuals ranged from -2.80 to 2.83 , indicating the absence of extreme outliers. Visual inspection of residual plots, including Q-Q plots and histograms, further suggested that the assumptions of normality and homoscedasticity were reasonably satisfied.

The nonparametric dispersion test revealed no evidence of over-dispersion or under-dispersion (dispersion = 0.97, $p = .808$), indicating that the variance structure of the model was appropriate for the present data. Additionally, the outlier test based on a binomial distribution revealed no significant outliers (0 outliers detected, $p = .13$), suggesting that no observations exerted undue influence on the model. Overall, these diagnostic analyses suggested that the linear mixed-effects regression model provided an acceptable fit to the secondary task data and was appropriate for examining the main and interaction effects of reverberation time, signal-to-noise ratio, and speaker group on listening effort performance.

In addition, the random effects structure indicated that the variance attributable to individual differences across participants was 0.44 (SD = 0.66), whereas the residual variance was 0.77 (SD = 0.88). This pattern suggests that although there was moderate between-subject variability, a larger proportion of the variance was explained by within-subject changes across experimental conditions. Thus, acoustic manipulations such as reverberation and SNR exerted a stronger influence on listening effort performance than inherent participant differences.

A LMER model was constructed with RT, SNR, and language group included as fixed effects, and Participant_ID specified as a random effect. The model is represented by equation 2:

$$\text{Secondary task scores} \sim \text{RT} * \text{SNR} * \text{Group} + (1 | \text{Participant_ID}) - \text{eq (2)}$$

The LMER model revealed significant fixed effects of RT, SNR, and language group on secondary task performance as shown in Table 4.6. With respect to the fixed effects, the intercept estimate ($\beta = 10.87$, $\text{SE} = 0.20$, $t = 53.97$, $p < .001$) represented the predicted secondary task performance under the baseline condition defined as minimum

reverberation, 0 dB SNR, and the Hindi speaker group. The correlation matrix of fixed effects was examined to assess multicollinearity among predictors. Correlations ranged from -0.71 to $+0.50$, indicating moderate associations among parameter estimates. Importantly, none of the correlations approached critical thresholds ($|r| \geq 0.80$), suggesting that multicollinearity was not a concern and that the parameter estimates were stable and reliable. The fixed-effects analysis revealed significant effects of RT, SNR, and language group, and their interactions on secondary task performance. Effect sizes were interpreted based on the magnitude of the standardized beta coefficients.

Table 4.6

Fixed-effect estimates from the LMER model showing the effects of group, RT, SNR, and on secondary task scores

Parameter	Estimate	Std. Error	df	<i>t</i> value	<i>p</i> value	Coefficient (β)
(Intercept)	10.87	0.20	249.33	53.97	< .001	-0.36
Max RT	-6.03	0.23	261.00	-26.55	< .001	-1.32
4 dB SNR	4.97	0.23	261.00	21.86	< .001	1.09
Malayalam	1.57	0.28	249.33	5.50	< .001	0.34
Kannada	4.40	0.28	249.33	15.45	< .001	0.96
Max RT $\times$ 4 dB SNR	-0.20	0.32	261.00	-0.62	0.53	N/A

Max RT × Malayalam	0.83	0.32	261.00	2.59	0.010	0.18
Max RT × Kannada	0.67	0.32	261.00	2.08	0.039	0.15
4 dB SNR × Malayalam	0.73	0.32	261.00	2.28	0.023	0.16
4 dB SNR × Kannada	0.87	0.32	261.00	2.70	0.007	0.19
(Max RT × 4 dB SNR) × Malayalam	-1.97	0.45	261.00	-4.33	<	-0.43
Malayalam					.001	
(Max RT × 4 dB SNR) × Kannada	-1.30	0.45	261.00	-2.86	0.005	-0.29
Kannada						

Note: SNR -signal to noise ratio ; RT-reverberation time ; Min -minimum; Avg-average ; Max –maximum

Visual inspection of Table 4.6 revealed a very large negative effect of RT on secondary task performance, indicating that listening effort significantly increased as reverberation increased. Compared to the reference condition (Minimum RT), Maximum RT demonstrated a very large negative effect on secondary task scores ($\beta = -1.32$, 95% CI $[-1.42, -1.22]$, $t(346) = -26.55$, $p < .001$). Examination of effect size as indicated by β revealed a very large detrimental effect, implying that increased reverberation substantially reduced secondary task performance and increased cognitive listening effort across participants.

SNR also demonstrated a very large positive effect, with significantly better secondary task performance observed at 4 dB SNR compared to 0 dB SNR ($\beta = 1.09$, 95% CI $[0.99, 1.19]$, $t(346) = 21.86$, $p < .001$). Examination of effect size revealed a very large positive effect, suggesting that improved signal clarity considerably reduced listening effort and facilitated better dual-task performance.

Speaker group had a substantial influence on secondary task performance. Relative to the Hindi group, the Malayalam group demonstrated a small-to-moderate positive effect ($\beta = 0.34$, 95% CI [0.22, 0.47], $t(346) = 5.50$, $p < .001$), whereas the Kannada group demonstrated a large positive effect ($\beta = 0.96$, 95% CI [0.84, 1.09], $t(346) = 15.45$, $p < .001$), indicating superior secondary task performance among native Kannada speakers. Examination of effect size revealed that native Kannada speakers exhibited substantially lower listening effort compared to the non-native speaker groups, suggesting that linguistic familiarity facilitated more efficient cognitive processing during speech perception in noise.

Comparison of effect sizes revealed that RT exerted the strongest influence on secondary task performance, followed by SNR and then language group. The very large negative effect associated with Maximum RT indicated that reverberation imposed the greatest cognitive demand on speakers. SNR demonstrated the second strongest effect, indicating that improved acoustic clarity substantially reduced listening effort. Although speaker group effects were comparatively smaller, Kannada speakers consistently outperformed the non-native speaker groups, demonstrating the beneficial effect of native language familiarity on listening effort performance.

The interaction between Maximum RT and 4 dB SNR demonstrated a negligible and non-significant effect ($\beta = -0.04$, 95% CI [-0.18, 0.09], $t(346) = -0.62$, $p = .534$), suggesting that the influence of reverberation on listening effort did not substantially differ across SNR conditions. However, interactions between RT and language group revealed significant positive effects. Maximum RT $\times$ Malayalam demonstrated a small positive effect ($\beta = 0.18$, $p = .010$), whereas Maximum RT $\times$ Kannada demonstrated a small positive effect ($\beta = 0.15$, $p = .039$), indicating that the detrimental effect of reverberation varied across speaker groups.

Similarly, the interaction between 4 dB SNR and language group revealed small positive effects for both Malayalam ($\beta = 0.16, p = .023$) and Kannada speakers ($\beta = 0.19, p = .007$), suggesting that improvements in SNR benefited these speaker groups differently relative to the Hindi group. Examination of effect sizes indicated that linguistic background moderately influenced the extent to which speakers benefited from improved acoustic conditions.

Among the three-way interactions, Maximum RT $\times$ 4 dB SNR $\times$ Malayalam demonstrated a moderate negative effect ($\beta = -0.43$, 95% CI $[-0.63, -0.24]$, $t(346) = -4.33$, $p < .001$), while Maximum RT $\times$ 4 dB SNR $\times$ Kannada also showed a moderate negative effect ($\beta = -0.29$, 95% CI $[-0.48, -0.09]$, $t(346) = -2.86$, $p = .004$). These findings imply that the combined effects of reverberation and noise differentially influenced listening effort across speaker groups, with non-native speakers demonstrating greater susceptibility to acoustically adverse conditions compared to native Kannada speakers.

The main effects of RT, SNR, and speaker group on secondary task performance were examined using Type III Wald χ^2 tests obtained from the linear mixed-effects model. The analysis revealed highly significant main effects of RT, [$\chi^2(1) = 705.09, p < .001, \eta_p^2 = 0.94$], SNR, [$\chi^2(1) = 477.81, p < .001, \eta_p^2 = 0.91$], and speaker group, [$\chi^2(2) = 245.39, p < .001, \eta_p^2 = 0.87$], as shown in Table 4.7. Examination of effect sizes revealed very large effects for RT, SNR, and speaker group, indicating that all three factors strongly influenced secondary task performance and listening effort.

Table 4.7

ANOVA Results (Type II Wald χ^2 Tests) for Secondary Task Performance

Effect	χ^2	df	95% CI	p-value	Effect Size (η_p^2)
--------	----------	----	--------	---------	----------------------------

RT	705.09	1	[0.94, 1.00]	< .001	0.94
SNR	477.81	1	[0.90, 1.00]	< .001	0.91
Group	245.39	2	[0.83, 1.00]	< .001	0.87
RT × SNR	0.39	1	[0.09, 1.00]	0.534	N/A
RT × Group	7.53	2	[0.00, 1.00]	0.023	0.003
SNR × Group	8.44	2	[0.00, 1.00]	0.015	0.02
RT × SNR × Group	19.38	2	[0.02, 1.00]	< .001	0.07

Note: SNR -signal to noise ratio; RT-reverberation time

A significant negative effect of RT was observed, indicating that increasing reverberation adversely affected secondary task performance. As reverberation increased from Minimum RT to Maximum RT, secondary task scores declined substantially, suggesting that highly reverberant listening conditions increased cognitive listening effort and reduced participants' ability to efficiently perform the secondary task. The very large effect size associated with RT implies that reverberation exerted the strongest influence on listening effort among all the factors examined.

Similarly, a very large main effect of SNR was observed, demonstrating that participants performed significantly better under the higher SNR condition (4 dB SNR)

compared to the poorer SNR condition (0 dB SNR). This finding indicates that improved signal clarity reduced listening effort and facilitated more efficient allocation of cognitive resources during dual-task performance. The large effect size associated with SNR further suggests that acoustic clarity strongly influenced listening effort performance.

Language of the speaker also demonstrated a very large effect on secondary task performance. Native Kannada speakers consistently outperformed the non-native speaker groups, indicating reduced listening effort and more efficient speech processing among native speakers. The substantial effect size associated with speaker group suggests that linguistic familiarity played an important role in modulating cognitive effort during speech perception in adverse listening environments.

Further examination of interaction effects revealed a non-significant $RT \times SNR$ interaction, [$\chi^2(1) = 0.39, p = .534, \eta_p^2 = 0.16$]. Although the effect size suggested a moderate interaction effect, the non-significant result indicates that the influence of reverberation on listening effort did not vary substantially across SNR conditions.

The $RT \times Group$ interaction was statistically significant, [$\chi^2(2) = 7.53, p = .023, \eta_p^2 = 0.003$], suggesting that the influence of reverberation differed across speaker groups. However, examination of effect size revealed only a very small effect, indicating that although speaker groups responded differently to reverberation, the magnitude of this difference was limited.

Similarly, the $SNR \times Group$ interaction was statistically significant, [$\chi^2(2) = 8.44, p = .015, \eta_p^2 = 0.02$], indicating that the benefit obtained from improved SNR varied across speaker groups. Examination of effect size revealed a small effect, suggesting that linguistic background moderately influenced the extent to which speakers benefited from improved acoustic clarity.

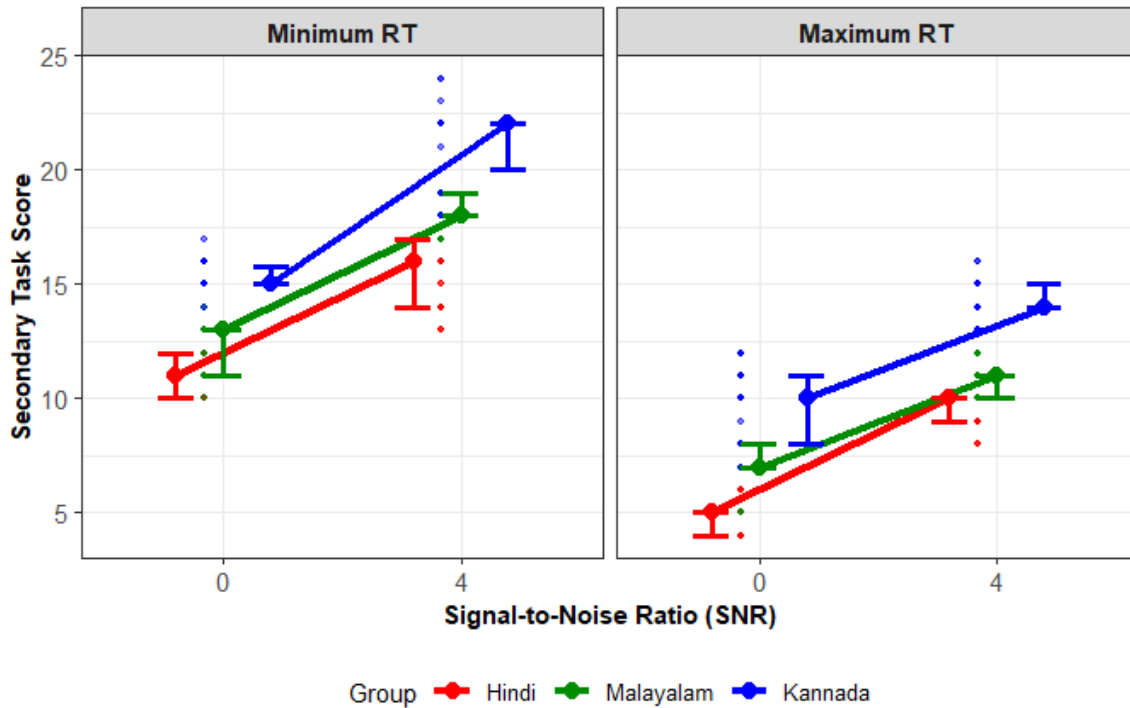
Importantly, the three-way interaction between RT, SNR, and speaker group was statistically significant, [$\chi^2(2) = 19.38, p < .001, \eta_p^2 = 0.07$]. Examination of effect size revealed a moderate effect, indicating that the combined effects of reverberation and noise on listening effort differed meaningfully across speaker groups. Specifically, non-native speakers demonstrated greater listening effort under acoustically adverse conditions compared to native Kannada speakers.

4.3.1. Effect of Language of speaker Group on Secondary Task Performance

As LMER model revealed significant main effect of Groups, as well as its significant interaction with SNR and RT (Table 4.6), pairwise comparison with Bonferroni correction were carried out for comparison of secondary task scores of three Groups across two SNRs and RT conditions. Figure 4.5 shows interaction plots with median and IQR along with individual data points for secondary task as a function of Group.

Figure 4.5

Interaction plots illustrating secondary task performance across SNR conditions under Minimum and Maximum RT conditions for Hindi, Malayalam, and Kannada speaker groups. Coloured lines represent the median secondary task scores for each speaker group across SNR conditions. Individual dots indicate each participant performance and Error bars represent the interquartile range (IQR).



Inspection of the Figure 4.5 reveals clear group-wise simple effects at each RT $\times$ SNR combination. Under minimum RT at 0 dB SNR, the Kannada group demonstrates the highest secondary task scores, followed by the Malayalam group, with the Hindi group showing the lowest performance. A similar pattern is observed at minimum RT at 4 dB SNR, where all groups show improved scores; however, the Kannada group continues to outperform the other groups, the Malayalam group remains intermediate, and the Hindi group shows comparatively lower scores. The magnitude of improvement from 0 dB to 4 dB appears evident across all groups, with a slightly steeper increase for the Kannada and Malayalam groups relative to the Hindi group.

Under maximum RT at 0 dB SNR, overall performance is reduced compared to minimum RT, but the same ordering of groups is maintained, with Kannada showing the highest scores, followed by Malayalam, and Hindi exhibiting the lowest scores. At maximum RT at 4 dB SNR, secondary task scores increase for all groups; nevertheless, the

Kannada group continues to show superior performance, Malayalam remains intermediate, and Hindi shows the lowest scores. Although SNR-related improvement is evident, the absolute scores remain lower than those observed under minimum RT.

Overall, at each $RT \times SNR$ combination, the Kannada group consistently exhibits higher secondary task performance, followed by the Malayalam group, with the Hindi group showing the lowest performance. This consistent ordering across conditions indicates that group-related differences in listening effort persist irrespective of acoustic condition, although all groups benefit from improved SNR and are adversely affected by increased reverberation.

Post hoc pairwise comparisons performed using estimated marginal means (EMMs) with Bonferroni adjustment for multiple comparisons, showed significant differences in secondary task scores between Group conditions as shown in Table 4.8. Pairwise comparisons for speaker group differences across RT and SNR conditions revealed significant differences between all groups under nearly all listening conditions. Examination of effect sizes indicated predominantly large to very large effects, suggesting substantial differences in listening effort performance across native and non-native speaker groups.

At Minimum RT and 0 dB SNR, the Hindi group performed significantly poorer than the Malayalam group (estimate = -1.57 , $t = -5.50$, $p < .001$), with a large effect size ($d = -1.78$). The comparison between the Hindi and Kannada groups demonstrated a very large effect (estimate = -4.40 , $t = -15.45$, $p < .001$, $d = -5.00$), while the Malayalam and Kannada groups also differed significantly with a very large effect size (estimate = -2.83 , $t = -9.95$, $p < .001$, $d = -3.22$). These findings imply that native Kannada speakers experienced substantially lower listening effort and better secondary task performance compared to non-native speaker groups under adverse noise conditions.

At Maximum RT and 0 dB SNR, all group comparisons remained statistically significant. Hindi speakers demonstrated significantly poorer performance compared to Malayalam speakers (estimate = -2.40 , $t = -8.43$, $p < .001$), with a large effect size ($d = -2.73$). The Hindi versus Kannada comparison showed a very large effect (estimate = -5.07 , $t = -17.79$, $p < .001$, $d = -5.76$), while Malayalam versus Kannada also demonstrated a very large effect size (estimate = -2.67 , $t = -9.37$, $p < .001$, $d = -3.03$). These results suggest that increasing reverberation further amplified listening effort demands, particularly among non-native speakers.

At Minimum RT and 4 dB SNR, significant differences were again observed across all speaker groups. Hindi speakers performed significantly poorer than Malayalam speakers (estimate = -2.30 , $t = -8.08$, $p < .001$), with a large effect size ($d = -2.61$). Comparisons between Hindi and Kannada speakers revealed a very large effect (estimate = -5.27 , $t = -18.50$, $p < .001$, $d = -5.98$), while Malayalam versus Kannada speakers also showed a very large effect size (estimate = -2.97 , $t = -10.42$, $p < .001$, $d = -3.37$). These findings indicate that even under improved SNR conditions, native Kannada speakers maintained a substantial advantage in listening effort performance.

At Maximum RT and 4 dB SNR, all group comparisons continued to show significant differences. The Hindi versus Malayalam comparison demonstrated a moderate-to-large effect (estimate = -1.17 , $t = -4.10$, $p = .0002$, $d = -1.33$). The Hindi versus Kannada comparison showed a very large effect size (estimate = -4.63 , $t = -16.27$, $p < .001$, $d = -5.27$), while Malayalam versus Kannada speakers also differed significantly with a very large effect (estimate = -3.47 , $t = -12.18$, $p < .001$, $d = -3.94$). These findings imply that the combined effects of reverberation and noise disproportionately increased listening effort among non-native speakers, whereas native Kannada speakers demonstrated greater resilience under acoustically challenging conditions.

Overall comparison of effect sizes across speaker groups revealed that the largest differences in secondary task performance were consistently observed between the Hindi and Kannada groups, followed by the Malayalam and Kannada groups, while the smallest differences were observed between the Hindi and Malayalam groups. The Hindi versus Kannada comparisons demonstrated very large effect sizes across nearly all RT and SNR conditions, indicating substantial differences in listening effort performance between native and non-native speakers. Malayalam versus Kannada comparisons also showed large to very large effects, although the magnitude was comparatively smaller than the Hindi–Kannada comparisons. In contrast, the Hindi versus Malayalam comparisons demonstrated relatively smaller, though still moderate-to-large effects. These findings imply that native Kannada speakers consistently experienced lower listening effort and superior secondary task performance compared to both non-native speaker groups, whereas the Hindi and Malayalam groups demonstrated comparatively similar levels of listening effort under adverse listening conditions.

Table 4.8

Pairwise Group Comparisons on Secondary Task Performance Across RT and SNR Conditions

SNR	RT	Hindi vs Malayalam				Hindi vs Kannada				Malayalam vs Kannada			
		Estimate	SE	<i>t</i> value,	Effect	Estimate	SE	<i>t</i> value,	Effect	Estimate	SE	<i>t</i> value,	Effect
				<i>p</i> value	Size			<i>p</i> value	Size			<i>p</i> value	Size
					(<i>d</i>)				(<i>d</i>)				(<i>d</i>)
0 dB	Min	-1.57	0.29	-5.50, < .001	-1.78	-4.40	0.29	-15.45, < .001	-5.00	-2.83	0.29	-9.95, < .001	-3.22
0 dB	Max	-2.40	0.29	-8.43, < .001	-2.73	-5.07	0.29	-17.79, < .001	-5.76	-2.67	0.29	-9.37, < .001	-3.03
4 dB	Min	-2.30	0.29	-8.08, < .001	-2.61	-5.27	0.29	-18.50, < .001	-5.98	-2.97	0.29	-10.42, < .001	-3.37
4 dB	Max	-1.17	0.29	-4.10, .0002	-1.33	-4.63	0.29	-16.27, < .001	-5.27	-3.47	0.29	-12.18, < .001	-3.94

Note: SNR -signal to noise ratio ; RT-reverberation time ; Min -minimum ; Avg-average ; Max -maximum

Interaction Effects Involving Group

The interaction between RT and Group was statistically significant, [$\chi^2(2) = 7.53$, $p = .023$, $\eta_p^2 = 0.003$], indicating that the effect of reverberation differed across speaker groups. Examination of effect size revealed a very small interaction effect, suggesting that although all speaker groups demonstrated increased listening effort with increasing reverberation, the magnitude of this effect varied slightly across native and non-native speakers. Non-native speakers demonstrated comparatively greater reductions in secondary task performance under highly reverberant conditions compared to native Kannada speakers.

Similarly, the SNR $\times$ Group interaction was statistically significant, [$\chi^2(2) = 8.44$, $p = .015$, $\eta_p^2 = 0.02$], suggesting that the speaker groups differed in the extent to which they benefited from improved SNR conditions. Examination of effect size revealed a small interaction effect, indicating that linguistic background moderately influenced the degree to which speakers benefited from improved acoustic clarity. Native Kannada speakers consistently demonstrated superior secondary task performance and lower listening effort across SNR conditions.

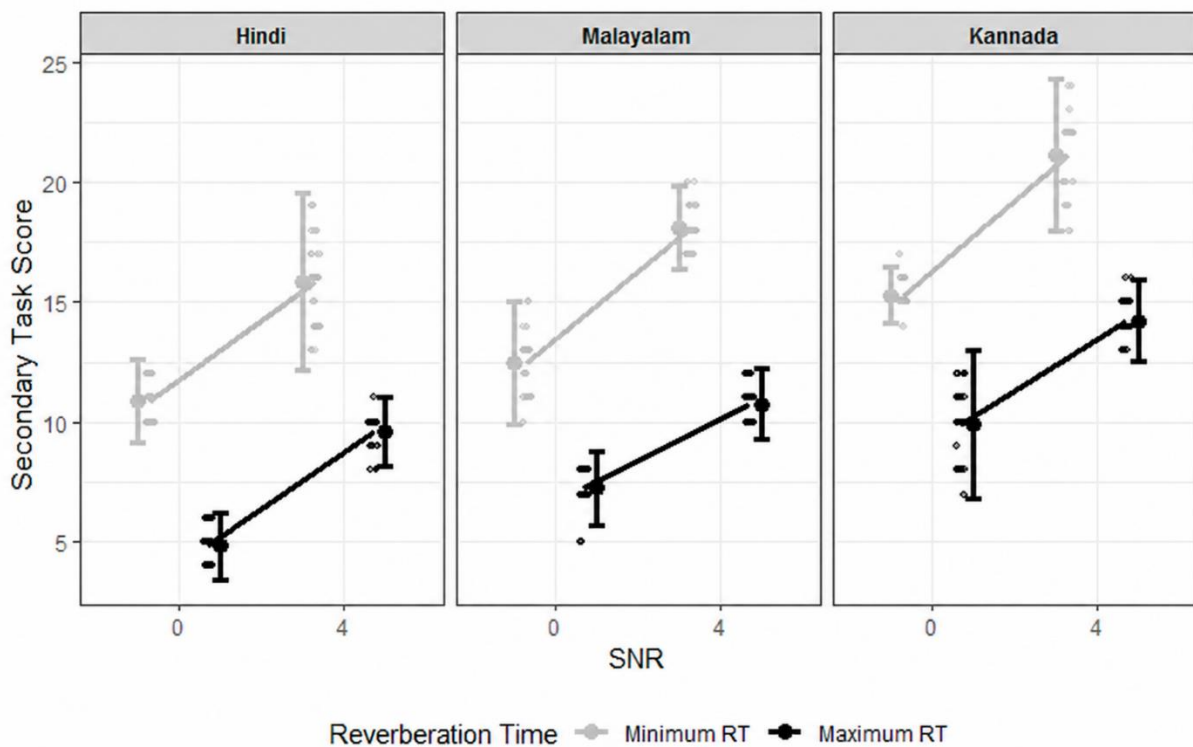
Importantly, the three-way interaction between RT, SNR, and Group was statistically significant, [$\chi^2(2) = 19.38$, $p < .001$, $\eta_p^2 = 0.07$]. Examination of effect size revealed a moderate interaction effect, indicating that the combined effects of reverberation and noise varied meaningfully as a function of linguistic background. Specifically, non-native speakers demonstrated greater listening effort and reduced secondary task performance under acoustically adverse conditions, whereas native Kannada speakers showed comparatively greater resilience to the combined degradations of reverberation and background noise.

4.3.2 Effect of Reverberation (RT) on Secondary Task Performance

As LMER model revealed significant main effect of RT as well as its significant interaction with Group and SNR (Table 4.6), pairwise comparison with Bonferroni correction were carried out for comparison of Secondary Task scores of two RT conditions for each group across SNRs. Figure 4.6 shows interaction plots with median and IQR along with individual data points for secondary task as a function of RT.

Figure 4.6

Interaction plots illustrating secondary task performance across SNR conditions for the Hindi, Malayalam, and Kannada speaker groups under Minimum and Maximum conditions. Coloured lines represent the median secondary task scores across SNR conditions for each reverberation condition. Individual dots indicate each participant performance and Error bars represent the interquartile range (IQR).



Inspection of the Figure 4.6 demonstrates clear RT-wise simple effects at each SNR $\times$ Group combination. At 0 dB SNR for the Hindi group, secondary task scores are higher under minimum RT compared to maximum RT, indicating better performance in less reverberant conditions. A similar pattern is observed for the Malayalam group at 0 dB SNR, where scores under minimum RT exceed those under maximum RT. For the Kannada group at 0 dB SNR, secondary task performance is also notably higher under minimum RT than maximum RT, with a clear separation between the two conditions.

At 4 dB SNR for the Hindi group, secondary task scores again show higher values under minimum RT compared to maximum RT, reflecting reduced listening effort in more favorable acoustic conditions. The Malayalam group at 4 dB SNR follows the same trend, with improved performance under minimum RT relative to maximum RT. Similarly, for the Kannada group at 4 dB SNR, scores under minimum RT remain higher than those under maximum RT, with the largest absolute scores observed in this condition.

Overall, across all SNR $\times$ Group combinations, secondary task performance is consistently higher under minimum RT than maximum RT, indicating that increased reverberation negatively impacts performance and elevates listening effort. Although performance improves with increasing SNR within each RT condition, the advantage of minimum RT persists across all groups, suggesting that reduced reverberation facilitates more efficient cognitive processing during listening tasks.

Pairwise comparisons of RT conditions for the secondary task revealed significant differences between Minimum RT and Maximum RT across all speaker groups and SNR conditions, as shown in Table 4.9. Examination of effect sizes (d) revealed consistently very

large effects, indicating that increasing reverberation substantially increased listening effort and reduced secondary task performance across participants.

At 0 dB SNR, the Hindi group demonstrated a very large effect for the comparison between Minimum RT and Maximum RT (estimate = 5.70, $t = 27.91$, $p < .001$, $d = 4.80$). Similarly, the Malayalam group showed a very large effect (estimate = 5.20, $t = 26.71$, $p < .001$, $d = 4.60$), while the Kannada group also demonstrated a very large effect (estimate = 4.30, $t = 23.80$, $p < .001$, $d = 4.10$). These findings imply that increased reverberation significantly increased listening effort under poorer SNR conditions for all speaker groups.

At 4 dB SNR, significant differences between Minimum RT and Maximum RT were again observed across all groups. The Hindi group demonstrated a very large effect (estimate = 4.10, $t = 20.56$, $p < .001$, $d = 3.50$), followed by the Malayalam group (estimate = 3.90, $t = 20.15$, $p < .001$, $d = 3.40$), and the Kannada group (estimate = 3.20, $t = 19.49$, $p < .001$, $d = 3.00$). These results indicate that even under improved acoustic clarity, increasing reverberation continued to significantly elevate listening effort and reduce secondary task performance.

Overall comparison of effect sizes revealed that the Hindi group demonstrated the largest RT-related effects, followed by the Malayalam group, while the Kannada group demonstrated comparatively smaller, though still very large, effects across both SNR conditions. This pattern suggests that non-native speakers were more adversely affected by increasing reverberation than native Kannada speakers. Furthermore, effect sizes were consistently larger under 0 dB SNR compared to 4 dB SNR, implying that the detrimental influence of reverberation on listening effort became more pronounced under poorer acoustic conditions.

Table 4.9*Pairwise Comparison of RT Conditions on Secondary Task Performance*

RT	SNR	Hindi				Malayalam				Kannada			
		Estimate	SE	<i>t, p</i>	Effect Size (<i>d</i>)	Estimate	SE	<i>t, p</i>	Effect Size (<i>d</i>)	Estimate	SE	<i>t, p</i>	Effect Size (<i>d</i>)
Min	vs 0 dB	5.70	0.21	27.91,	4.80	5.20	0.19	26.71,	4.60	4.30	0.18	23.80,	4.10
Max				< .001				< .001				< .001	
	4 dB	4.10	0.20	20.56,	3.50	3.90	0.19	20.15,	3.40	3.20	0.16	19.49,	3.00
				< .001				< .001				< .001	

Note: SNR -signal to noise ratio; RT-reverberation time; Min -minimum; Avg-average; Max -maximum

Interaction Effects Involving RT

The interaction between RT and SNR was not statistically significant, [$\chi^2(1) = 0.39$, $p = .534$, $\eta_p^2 = 0.16$], indicating that the effect of reverberation on secondary task performance did not differ substantially across SNR conditions. However, examination of effect size revealed a moderate interaction effect, suggesting that reverberation may still have exerted some influence on listening effort across different SNR conditions, although the effect was not statistically reliable.

The RT $\times$ Group interaction was statistically significant, [$\chi^2(2) = 7.53$, $p = .023$, $\eta_p^2 = 0.003$], suggesting that reverberation affected the speaker groups differently. Although all groups demonstrated reduced secondary task performance with increasing RT, the non-native speaker groups showed comparatively greater increases in listening effort under Maximum RT conditions. Examination of effect size revealed a very small interaction effect, indicating that the magnitude of group differences associated with reverberation was relatively limited.

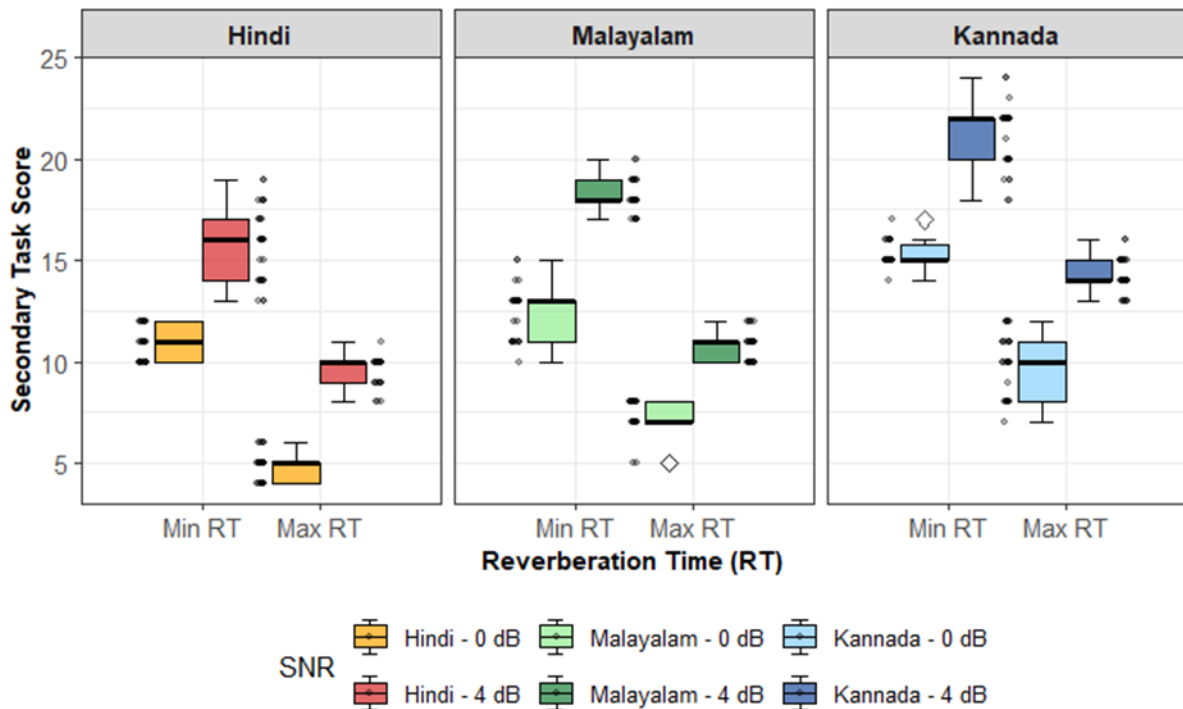
Further, the three-way interaction between RT, SNR, and Group was statistically significant, [$\chi^2(2) = 19.38$, $p < .001$, $\eta_p^2 = 0.07$], indicating that the combined effects of reverberation and noise differed across speaker groups. Non-native speakers demonstrated greater listening effort and poorer secondary task performance under acoustically adverse conditions. Examination of effect size revealed a moderate interaction effect, suggesting that linguistic background meaningfully influenced the combined effects of reverberation and noise on listening effort performance.

4.3.3 Effect of Signal-to-Noise Ratio (SNR) on Secondary Task Performance

As LMER model revealed significant main effect of SNR, as well as its significant interaction with Group and RT (Table 4.6), pairwise comparison with Bonferroni correction were carried out for comparison of Secondary Task scores of two SNR conditions for each group across RTs. Figure 4.7 shows boxplots with median and IQR along with individual data points for secondary task as a function of SNR.

Figure 4.7

Boxplots illustrating secondary task performance across RT conditions (Minimum and Maximum RT) for Hindi, Malayalam, and Kannada speaker groups under 0 dB and 4 dB SNR conditions. The central horizontal line within each box represents the median score, while the box boundaries indicate the interquartile range (IQR). Whiskers represent the spread of the data beyond the quartiles, and individual dots indicate each participant performance and outliers.



Inspection of Figure 4.7 revealed clear SNR-wise simple effects at each $RT \times Group$ combination. Under minimum RT for the Hindi group, secondary task scores are higher at 4 dB SNR compared to 0 dB SNR, indicating improved performance with better listening conditions. A similar pattern is observed for the Malayalam group under minimum RT, where scores at 4 dB SNR exceed those at 0 dB SNR, reflecting reduced listening effort. For the Kannada group under minimum RT, the increase from 0 dB to 4 dB SNR is also evident, with the highest overall scores observed at 4 dB SNR among all groups.

Under maximum RT for the Hindi group, secondary task scores again show an increase at 4 dB SNR relative to 0 dB SNR, although the overall scores are lower compared to minimum RT. The Malayalam group under maximum RT demonstrates a similar improvement with increasing SNR, with higher scores at 4 dB SNR than at 0 dB SNR. Likewise, for the Kannada group under maximum RT, secondary task performance improves at 4 dB SNR, maintaining the same directional trend despite reduced overall performance due to reverberation.

Overall, at each $RT \times Group$ combination, secondary task scores are consistently higher at 4 dB SNR than at 0 dB SNR, indicating that improved signal-to-noise conditions reduce listening effort across all speaker groups. However, the magnitude of improvement appears more pronounced under minimum RT compared to maximum RT, suggesting that the beneficial effect of SNR is somewhat constrained in the presence of higher reverberation.

To further examine the effect of SNR, pairwise comparisons across groups at each $RT \times SNR$ condition was conducted using Bonferroni correction. Pairwise comparisons of

SNR conditions for the secondary task revealed significant differences between 0 dB SNR and 4 dB SNR across all RT and speaker group conditions, as shown in Table 4.10. Examination of effect sizes revealed predominantly large to very large effects, indicating that improvements in acoustic clarity substantially reduced listening effort and enhanced secondary task performance.

Under Minimum RT conditions, the Hindi group demonstrated a very large effect for the comparison between 0 dB SNR and 4 dB SNR (estimate = -3.10 , $t = -14.86$, $p < .001$, effect size (d) = -2.70). Similarly, the Malayalam group demonstrated a very large effect (estimate = -2.90 , $t = -11.23$, $p < .001$, $d = -2.30$), while the Kannada group showed a large effect (estimate = -1.40 , $t = -3.77$, $p = .003$, $d = -0.90$). These findings imply that improved signal clarity significantly reduced listening effort for all speaker groups, although the magnitude of improvement was comparatively smaller for native Kannada speakers.

Under Maximum RT conditions, all speaker groups again demonstrated significant differences between SNR conditions. The Hindi group showed a very large effect (estimate = -5.20 , $t = -28.69$, $p < .001$, $d = -4.90$), followed by the Malayalam group (estimate = -4.10 , $t = -17.08$, $p < .001$, $d = -3.60$), and the Kannada group (estimate = -2.80 , $t = -12.99$, $p < .001$, $d = -2.10$). These findings indicate that improvements in SNR had an even greater influence under highly reverberant conditions, substantially reducing listening effort and improving secondary task performance.

Overall comparison of effect sizes revealed that the Hindi group demonstrated the largest SNR-related effects, followed by the Malayalam group, whereas the Kannada group demonstrated comparatively smaller effects across both RT conditions. This pattern suggests that non-native speakers derived greater benefit from improved acoustic clarity than native Kannada speakers. Furthermore, effect sizes were consistently larger under Maximum RT conditions compared to Minimum RT conditions, implying that improved SNR became

increasingly important for reducing listening effort under highly reverberant listening environments.

Table 4.10

Pairwise Comparison of SNR Conditions (0 dB vs 4 dB) on Secondary Task Performance

SNR	RT	Hindi				Malayalam				Kannada			
		Estimate	SE	<i>t, p</i>	Effect Size (<i>d</i>)	Estimate	SE	<i>t, p</i>	Effect Size (<i>d</i>)	Estimate	SE	<i>t, p</i>	Effect Size (<i>d</i>)
0 dB vs 4 dB	Min	-3.10	0.18	-14.86,	-2.70	-2.90	0.19	-11.23,	-2.30	-1.40	0.17	-3.77,	-0.90
				< .001				< .001				.01	
dB	Max	-5.20	0.20	-28.69,	-4.90	-4.10	0.19	-17.08,	-3.60	-2.80	0.18	-12.99,	-2.10
				< .001				< .001				< .001	

Note: SNR -signal to noise ratio ; RT-reverberation time; Min- minimum; Max-maximum

Interaction Effects Involving SNR

The interaction between RT and SNR was not statistically significant, [$\chi^2(1) = 0.39$, $p = .534$, $\eta_p^2 = 0.16$], indicating that the effect of SNR on secondary task performance did not vary substantially across reverberation conditions. However, examination of effect size revealed a moderate interaction effect, suggesting that improvements in acoustic clarity may still have influenced listening effort differently across RT conditions, although the effect was not statistically reliable.

The SNR $\times$ Group interaction was statistically significant, [$\chi^2(2) = 8.44$, $p = .015$, $\eta_p^2 = 0.02$], suggesting that the speaker groups differed in the extent to which they benefited from improved SNR conditions. Examination of effect size revealed a small interaction effect, indicating that linguistic background moderately influenced the degree to which speakers benefited from improved acoustic clarity. Native Kannada speakers consistently demonstrated lower listening effort and superior secondary task performance across SNR conditions compared to the non-native speaker groups.

Further, the three-way interaction between RT, SNR, and Group was statistically significant, [$\chi^2(2) = 19.38$, $p < .001$, $\eta_p^2 = 0.07$], indicating that the combined effects of reverberation and noise differed across speaker groups. Examination of effect size revealed a moderate interaction effect, suggesting that linguistic background meaningfully influenced the combined effects of reverberation and noise on listening effort performance. Specifically, non-native speakers demonstrated greater listening effort and poorer secondary task performance under acoustically adverse conditions, whereas native Kannada speakers showed comparatively greater resilience to the combined degradations of reverberation and background noise.

4.4 Effect of RT, SNR and Group on Subjective Listening Effort (NASA-TLX)

Subjective listening effort was assessed using the NASA Task Load Index (NASA-TLX) across varying SNR and RT conditions in native and non-native Kannada speakers. The NASA-TLX scores obtained across conditions were analyzed using a Linear Mixed-Effects Regression (LMER) model, with RT, SNR, and Group included as fixed effects and Participant ID included as a random effect. Prior to examining the effects of RT, SNR, and language group on NASA-TLX, diagnostic analyses were conducted to evaluate the adequacy of the model.

Residual diagnostics for the NASA-TLX were evaluated using the DHARMA package to assess the adequacy of the linear mixed-effects regression model. Visual inspection of the DHARMA residual plots demonstrated generally acceptable residual behaviour, although slight deviations from residual uniformity were observed. The Kolmogorov–Smirnov test for uniformity was statistically significant ($D = 0.113, p < .001$), indicating a modest deviation from the expected residual distribution. However, the deviation was relatively small and the visual inspection of residual plots, histograms, and Q-Q plots suggested that the residuals reasonably approximated the expected distribution.

Further diagnostic testing revealed that the dispersion test was non-significant (dispersion = 0.98, $p = .816$), indicating no evidence of over-dispersion or under-dispersion within the model. Similarly, the outlier test was non-significant ($p = .906$), suggesting that the observed number of outliers did not substantially exceed the expected frequency. These findings indicated that the residual deviations were limited and did not severely violate model assumptions.

Due to the slight deviation identified by the Kolmogorov–Smirnov test, alternative modelling approaches including Gamma and Beta regression models were also explored.

However, these alternative models demonstrated comparatively poorer model fit, with greater evidence of residual deviation, dispersion, and outlier presence than the LMER model. Therefore, the linear mixed-effects regression model was retained as the final model for analysis.

To further account for the modest residual deviations and improve the robustness of statistical inference, robust standard errors were computed using the *clubSandwich* package with CR2 correction factors (`coef_test(model_lmm, vcov = "CR2")`). The inclusion of robust standard errors provided more reliable parameter estimates and significance testing despite minor departures from residual assumptions.

Model fit indices including Akaike Information Criterion (AIC) and Bayesian Information Criterion (BIC) were additionally examined to compare competing models. The retained LMER model demonstrated comparatively better fit (i.e. lower values of AIC, BIC) and greater model stability relative to the alternative Gamma and Beta regression models, further supporting its suitability for analysing secondary task performance. Appendix G gives the details of AIC and BIC across models.

A LMER model was constructed with RT, SNR, and language group included as fixed effects, and Participant_ID specified as a random effect. The model is represented by equation 3:

$$\text{NASA task scores} \sim \text{RT} * \text{SNR} * \text{Group} + (1 | \text{Participant_ID}) - \text{eq (2)}$$

The LMER model revealed fixed effects of RT, SNR, and speaker group on NASA-TLX scores, as shown in Table 4.11. The fixed-effects analysis demonstrated significant main effects of SNR and speaker group, along with selected interaction effects, on subjective listening effort performance. However, the main effects of RT were statistically non-significant. Effect sizes were interpreted based on the magnitude of the standardized beta coefficients (β).

Table 4.11

Fixed-effect estimates from the LMER model showing the effects of group, RT, SNR, and on NASA-TLX scores.

Parameter	Estimate	SE	95% CI [Lower bound, Upper bound]	<i>t</i> value	<i>p</i> value	Coefficient (β)
(Intercept)	27.43	0.12	[0.27, 0.74]	4.22	< .001	0.51
Avg RT	0.26	0.10	[-0.14, 0.24]	0.55	0.581	N/A
Max RT	0.46	0.08	[-0.06, 0.25]	1.22	0.221	N/A
4 dB SNR	-0.79	0.07	[-0.30, -0.03]	-2.36	0.019	-0.16
Malayalam	-2.66	0.17	[-0.89, -0.22]	-3.24	0.001	-0.55
Kannada	-6.16	0.17	[-1.61, -0.94]	-7.51	< .001	-1.28
Avg RT $\times$ 4 dB SNR	0.94	0.13	[-0.05, 0.44]	1.55	0.122	N/A
Max RT $\times$ 4 dB SNR	0.70	0.10	[-0.06, 0.35]	1.40	0.162	N/A
SNR						
Avg RT $\times$ 0.14	0.14	0.14	[-0.24, 0.30]	0.22	0.828	N/A
Malayalam						
Max RT $\times$ 1.50	1.50	0.11	[0.10, 0.53]	2.85	0.004	0.31
Malayalam						
Avg RT $\times$ Kannada	0.59	0.14	[-0.14, 0.39]	0.90	0.367	N/A
Max RT $\times$ Kannada	1.43	0.11	[0.08, 0.51]	2.70	0.007	0.30
4 dB SNR $\times$ 0.45	0.45	0.10	[-0.10, 0.29]	0.95	0.340	N/A
Malayalam						
4 dB SNR $\times$ 0.05	0.05	0.10	[-0.18, 0.20]	0.10	0.918	N/A
Kannada						

(Avg RT × 4 dB SNR) × Malayalam	0.88	0.18	[-0.17, 0.53]	1.03	0.305	N/A
(Max RT × 4 dB SNR) Malayalam	-0.74	0.15	[-0.44, 0.14]	-1.04	0.301	N/A
(Avg RT × 4 dB SNR) × Kannada	-0.23	0.18	[-0.40, 0.30]	-0.26	0.793	N/A
(Max RT × 4 dB SNR) × Kannada	-0.07	0.15	[-0.30, 0.28]	-0.10	0.923	N/A

Note: SNR -signal to noise ratio ; RT-reverberation time ; Min -minimum ; Avg-average ; Max –maximum

Visual inspection of Table 4.11 revealed relatively small effects of RT on the subjective listening effort scores. Compared to the reference condition (Minimum RT), both Average RT ($\beta = 0.05$, 95% CI $[-0.14, 0.24]$, $t = 0.55$, $p = .581$) and Maximum RT ($\beta = 0.09$, 95% CI $[-0.06, 0.25]$, $t = 1.22$, $p = .221$) demonstrated small positive effects that were statistically non-significant. Examination of effect size as indicated by β revealed negligible to small effects for both Average RT and Maximum RT conditions, implying that reverberation alone did not substantially influence subjective listening effort ratings measured using NASA-TLX scores.

SNR demonstrated a small negative effect, with significantly lower NASA-TLX scores observed at 4 dB SNR compared to 0 dB SNR ($\beta = -0.16$, 95% CI $[-0.30, -0.03]$, $t = -2.36$, $p = .019$). Examination of effect size revealed a small effect, implying that improved acoustic clarity modestly reduced perceived listening effort and task demand among participants.

Speaker group demonstrated comparatively stronger effects on subjective listening effort. Relative to the Hindi group, the Malayalam group showed a moderate negative effect

($\beta = -0.55$, 95% CI $[-0.89, -0.22]$, $t = -3.24$, $p = .001$), whereas the Kannada group demonstrated a large negative effect ($\beta = -1.28$, 95% CI $[-1.61, -0.94]$, $t = -7.51$, $p < .001$). Examination of effect sizes revealed that native Kannada speakers consistently reported lower subjective listening effort compared to non-native speaker groups, implying that linguistic familiarity reduced the perceived cognitive demand associated with speech perception tasks.

The interaction between RT and SNR demonstrated small positive effects for both Average RT $\times$ 4 dB SNR ($\beta = 0.20$, $p = .122$) and Maximum RT $\times$ 4 dB SNR ($\beta = 0.15$, $p = .162$); however, these interactions were statistically non-significant. These findings imply that the influence of reverberation on subjective listening effort did not substantially differ across SNR conditions.

Interactions between RT and speaker group revealed differential effects across groups. Average RT $\times$ Malayalam demonstrated a negligible effect ($\beta = 0.03$, $p = .828$), whereas Maximum RT $\times$ Malayalam showed a moderate positive effect ($\beta = 0.31$, $p = .004$). Similarly, Average RT $\times$ Kannada demonstrated a small effect ($\beta = 0.12$, $p = .367$), while Maximum RT $\times$ Kannada demonstrated a moderate positive effect ($\beta = 0.30$, $p = .007$). These findings suggest that increases in reverberation under Maximum RT conditions produced relatively greater increases in subjective listening effort among Malayalam and Kannada speakers.

The interactions between 4 dB SNR and speaker groups demonstrated negligible to small effects. The interaction between 4 dB SNR and Malayalam showed a small positive effect ($\beta = 0.09$, $p = .340$), whereas 4 dB SNR $\times$ Kannada demonstrated a negligible effect ($\beta = 0.01$, $p = .918$). These findings imply that improvements in acoustic clarity produced relatively similar reductions in subjective listening effort across speaker groups.

Among the three-way interactions, Average RT $\times$ 4 dB SNR $\times$ Malayalam demonstrated a small positive effect ($\beta = 0.18, p = .305$), whereas Maximum RT $\times$ 4 dB SNR $\times$ Malayalam showed a small negative effect ($\beta = -0.15, p = .301$). Similarly, Average RT $\times$ 4 dB SNR $\times$ Kannada demonstrated a negligible negative effect ($\beta = -0.05, p = .793$), while Maximum RT $\times$ 4 dB SNR $\times$ Kannada also showed a negligible negative effect ($\beta = -0.01, p = .923$). However, all three-way interactions were statistically non-significant, suggesting that the combined effects of reverberation and noise on subjective listening effort did not differ substantially across speaker groups.

Comparison of effect sizes revealed that speaker group exerted the strongest influence on NASA-TLX scores, followed by SNR, whereas RT demonstrated the smallest effects overall. The Kannada group demonstrated the largest effect size in reducing subjective listening effort, followed by the Malayalam group, indicating that native speakers consistently perceived lower listening effort compared to non-native speakers. In contrast, reverberation demonstrated only negligible to small effects on subjective listening effort ratings, suggesting that participants perceived workload was influenced more strongly by linguistic familiarity and acoustic clarity than by reverberation alone.

The main effects of RT, SNR, and speaker group on NASA-TLX scores were examined using Type III Wald chi-square tests obtained from the LMER model, as shown in Table 4.12. The analysis revealed a non-significant main effect of RT [$\chi^2(2) = 1.51, p = .469, \eta_p^2 = 0.06$], indicating that reverberation alone did not significantly influence subjective listening effort ratings. Examination of effect size revealed a small effect, implying that increases in reverberation produced only minimal changes in perceived workload and listening effort.

A statistically significant main effect of SNR was observed [$\chi^2(1) = 5.55, p = .019, \eta_p^2 = 2.13 \times 10^{-4}$], indicating that subjective listening effort differed across SNR conditions. Examination of effect size revealed a negligible effect, suggesting that improved acoustic clarity at 4 dB SNR only modestly reduced perceived listening effort among participants.

Language group of the speakers demonstrated a highly significant main effect on NASA-TLX scores [$\chi^2(2) = 56.76, p < .001, \eta_p^2 = 0.38$]. Examination of effect size revealed a large effect, indicating that linguistic background strongly influenced subjective listening effort ratings. Native Kannada speakers consistently reported lower listening effort compared to non-native speaker groups, implying that familiarity with the target language substantially reduced perceived cognitive demand during speech perception tasks.

The RT $\times$ SNR interaction was non-significant [$\chi^2(2) = 3.19, p = .203, \eta_p^2 = 5.10 \times 10^{-3}$]. Examination of effect size revealed a negligible interaction effect, suggesting that the influence of reverberation on subjective listening effort did not vary substantially across SNR conditions.

The RT $\times$ Group interaction was statistically significant [$\chi^2(4) = 11.18, p = .03, \eta_p^2 = 8.17 \times 10^{-3}$], indicating that reverberation affected speaker groups differently. Examination of effect size revealed a negligible to small interaction effect, implying that the effect of reverberation on subjective listening effort varied slightly across language groups, with non-native speakers demonstrating comparatively greater increases in perceived effort under adverse reverberation conditions.

The SNR $\times$ Group interaction was non-significant [$\chi^2(2) = 1.10, p = .58, \eta_p^2 = 1.54 \times 10^{-3}$]. Examination of effect size revealed a negligible effect, suggesting that improvements in SNR produced relatively similar reductions in subjective listening effort across speaker groups.

Similarly, the three-way interaction between RT, SNR, and Group was non-significant [$\chi^2(4) = 4.95$, $p = .29$, $\eta_p^2 = 2.30 \times 10^{-3}$]. Examination of effect size revealed a negligible interaction effect, indicating that the combined influence of reverberation and noise on subjective listening effort did not substantially differ across speaker groups.

Comparison of effect sizes revealed that speaker group exerted the strongest influence on NASA-TLX scores, followed by RT, whereas SNR demonstrated the smallest overall effect. The large effect associated with speaker group indicates that linguistic familiarity played the most important role in determining perceived listening effort. In contrast, RT and SNR demonstrated comparatively negligible to small effects, suggesting that subjective workload ratings were influenced more strongly by linguistic background than by acoustic degradation alone.

Table 4.12

ANOVA Results (Type II Wald χ^2 Tests) for NASA-TLX Performance

Effect	χ^2	df	95% CI	p value	Effect Size (η_p^2)
RT	1.51	2	[0.04, 1.00]	0.469	N/A
SNR	5.55	1	[0.00, 1.00]	0.019	2.13×10^{-4}
Group	56.76	2	[0.25, 1.00]	< .001	0.38
RT $\times$ SNR	3.19	2	[0.00, 1.00]	0.203	N/A
RT $\times$ Group	11.18	4	[0.00, 1.00]	0.025	8.17×10^{-3}
SNR $\times$ Group	1.10	2	[0.00, 1.00]	0.578	N/A
RT $\times$ SNR $\times$ Group	4.95	4	[0.00, 1.00]	0.292	N/A

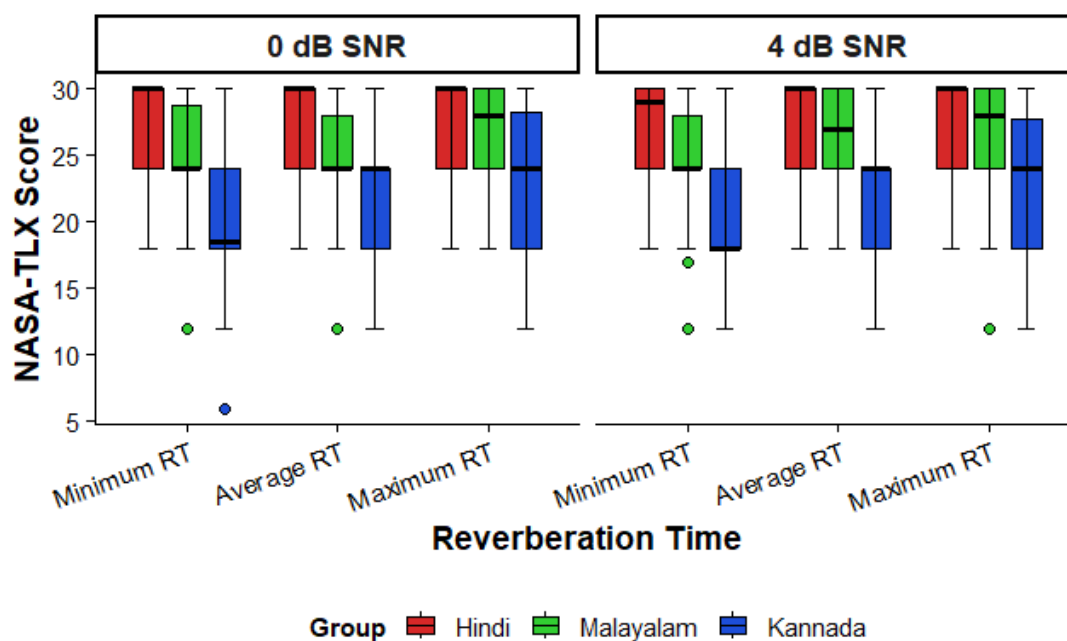
Note: SNR -signal to noise ratio ; RT-reverberation time

4.4.1 Effect of Language Group on NASA-TLX Scores

As LMER model revealed significant main effect of Groups, as well as its significant interaction with SNR and RT (Table 4.11), pairwise comparison with Bonferroni correction were carried out for comparison of NASA-TLX scores of three Groups across two SNRs and three RT conditions. Figure 4.8 shows boxplots with median and IQR along with individual data points for NASA-TLX as a function of Group.

Figure 4.8

Boxplots illustrating NASA-TLX performance across RT conditions (Minimum, Average and Maximum RT) for Hindi, Malayalam, and Kannada speaker groups under 0 dB and 4 dB SNR conditions. The central horizontal line within each box represents the median score, while the box boundaries indicate the interquartile range (IQR). Whiskers represent the spread of the data beyond the quartiles, and individual dots indicate each participant-performance and outliers.



Inspection of the figure 4.8 reveals clear group-wise differences in subjective listening effort across all acoustic conditions. Across both SNR conditions, Hindi-speaking non-native Kannada speakers consistently demonstrated the highest NASA-TLX scores across Minimum, Average, and Maximum RT conditions, indicating greater perceived listening effort. Malayalam-speaking speakers showed comparatively lower NASA-TLX scores than the Hindi group, whereas native Kannada speakers consistently exhibited the lowest listening effort across conditions.

Under the 0 dB SNR condition, subjective listening effort increased slightly with increasing reverberation time across all groups. The difference between groups was particularly evident under Average RT and Maximum RT conditions, where native Kannada speakers demonstrated noticeably lower NASA-TLX scores compared to the non-native speaker groups. Similar patterns were observed under the 4 dB SNR condition, although overall listening effort appeared slightly reduced relative to the 0 dB SNR condition.

The figure also demonstrates that the variability in NASA-TLX scores differed across groups. Kannada speakers showed relatively wider distributions and lower median scores under certain RT conditions, whereas Hindi speakers exhibited consistently elevated NASA-TLX scores with comparatively smaller variability. Malayalam speakers demonstrated intermediate performance between the two groups.

Overall, the graphical pattern suggests that linguistic familiarity with Kannada reduced perceived listening effort, while adverse acoustic conditions, particularly increased reverberation, contributed to greater subjective workload. These findings visually support

the significant main effect of speaker group and reverberation time observed in the mixed-effects analysis.

Post hoc pairwise comparisons performed using estimated marginal means (EMMs) with Bonferroni adjustment for multiple comparisons, showed significant differences in NASA-TLX scores between Group conditions as shown in Table 4.13. Pairwise comparisons of NASA-TLX scores across RT and SNR conditions revealed significant differences between speaker groups under several listening conditions, as shown in Table 4.13. Overall, Kannada speakers consistently demonstrated significantly lower NASA-TLX scores compared to both Hindi and Malayalam speakers, indicating reduced perceived listening effort among native Kannada speakers. Examination of effect sizes revealed predominantly large effects for comparisons involving Kannada speakers, suggesting that linguistic familiarity substantially reduced subjective listening effort.

Table 4.13

Pairwise Group Comparisons on NASA-TLX Performance Across RT and SNR Conditions

SNR	RT	Hindi vs Malayalam					Hindi vs Kannada					Malayalam vs Kannada				
		Estimate	SE	<i>t</i>	<i>p</i>	<i>d</i>	Estimate	SE	<i>t</i>	<i>p</i>	<i>d</i>	Estimate	SE	<i>t</i>	<i>p</i>	<i>d</i>
	Min	2.66	0.82	3.24	0.005	0.87	6.16	0.82	7.51	< .001	2.02	3.50	0.82	4.27	< .001	1.15
0 dB	Avg	2.52	0.93	2.72	0.022	0.83	5.57	0.93	6.01	< .001	1.83	3.05	0.93	3.29	0.004	1.00
	Max	1.16	0.84	1.38	0.509	N/A	4.73	0.84	5.64	< .001	1.55	3.58	0.84	4.26	< .001	1.17
	Min	2.21	0.81	2.73	0.022	0.72	6.11	0.81	7.57	< .001	2.01	3.91	0.81	4.84	< .001	1.28
4 dB	Avg	1.18	0.87	1.36	0.532	N/A	5.74	0.87	6.61	< .001	1.89	4.57	0.87	5.26	< .001	1.50
	Max	1.44	0.82	1.76	0.245	N/A	4.75	0.82	5.80	< .001	1.56	3.31	0.82	4.04	0.0003	1.09

Note: SNR -signal to noise ratio ; RT-reverberation time ; Min -minimum ; Avg-average ; Max -maximum

At 0 dB SNR under Minimum RT conditions, significant differences were observed between Hindi and Malayalam speakers ($t = 3.24, p = .005, d = 0.87$), Hindi and Kannada speakers ($t = 7.51, p < .001, d = 2.02$), and Malayalam and Kannada speakers ($t = 4.27, p < .001, d = 1.15$). Examination of effect sizes revealed a moderate effect for Hindi versus Malayalam and large effects for comparisons involving Kannada speakers, implying substantially lower listening effort among native Kannada speakers.

A similar pattern was observed at 0 dB SNR under Average RT conditions. Significant differences were obtained between Hindi and Malayalam speakers ($t = 2.72, p = .022, d = 0.83$), Hindi and Kannada speakers ($t = 6.01, p < .001, d = 1.83$), and Malayalam and Kannada speakers ($t = 3.29, p = .004, d = 1.00$). Examination of effect sizes again revealed moderate to large effects, indicating that Kannada speakers continued to experience lower subjective listening effort even under increased reverberation.

Under Maximum RT at 0 dB SNR, the comparison between Hindi and Malayalam speakers was non-significant ($t = 1.38, p = .509, d = 0.38$), demonstrating only a small effect. However, Hindi versus Kannada ($t = 5.64, p < .001, d = 1.55$) and Malayalam versus Kannada ($t = 4.26, p < .001, d = 1.17$) comparisons remained significant with large effect sizes, suggesting that Kannada speakers maintained lower perceived listening effort even under highly adverse acoustic conditions.

At 4 dB SNR under Minimum RT conditions, significant differences were observed between Hindi and Malayalam speakers ($t = 2.73, p = .022, d = 0.72$), Hindi and Kannada speakers ($t = 7.57, p < .001, d = 2.01$), and Malayalam and Kannada speakers ($t = 4.84, p < .001, d = 1.28$). Examination of effect sizes revealed moderate to very large effects, implying that improved acoustic clarity further reduced subjective listening effort among Kannada speakers.

At 4 dB SNR under Average RT conditions, the Hindi versus Malayalam comparison was non-significant ($t = 1.36, p = .532, d = 0.39$), indicating only a small effect. In contrast, Hindi versus Kannada ($t = 6.61, p < .001, d = 1.89$) and Malayalam versus Kannada ($t = 5.26, p < .001, d = 1.50$) comparisons demonstrated large effect sizes, suggesting substantially lower NASA-TLX scores among Kannada speakers.

Similarly, under Maximum RT at 4 dB SNR, Hindi versus Malayalam comparisons remained non-significant ($t = 1.76, p = .245, d = 0.47$), while Hindi versus Kannada ($t = 5.80, p < .001, d = 1.56$) and Malayalam versus Kannada ($t = 4.04, p = .0003, d = 1.09$) comparisons remained statistically significant with large effect sizes. These findings indicate that native Kannada speakers consistently experienced lower subjective listening effort across all reverberation and SNR conditions.

Comparison of effect sizes across conditions revealed that the largest effects were consistently observed for Hindi versus Kannada comparisons, followed by Malayalam versus Kannada comparisons, whereas Hindi versus Malayalam comparisons demonstrated comparatively smaller effects. Overall, Kannada speakers exhibited the lowest subjective listening effort ratings, followed by Malayalam speakers, while Hindi speakers demonstrated the highest perceived listening effort. These findings imply that linguistic familiarity played a major role in reducing subjective listening effort across varying acoustic conditions.

Interaction Effects Involving Group

The interaction between RT and Group was statistically significant, [$\chi^2(4) = 11.18, p = .025, \eta_p^2 = 8.17 \times 10^{-3}$], indicating that the effect of reverberation differed across speaker groups. Examination of effect size revealed a negligible to small interaction effect, suggesting that increases in reverberation produced relatively greater increases in subjective

listening effort among non-native speakers compared to native Kannada speakers. Although all speaker groups demonstrated changes in NASA-TLX scores across reverberation conditions, non-native speakers tended to report comparatively higher listening effort under adverse reverberation conditions.

Similarly, the SNR $\times$ Group interaction was non-significant [$\chi^2(2) = 1.10, p = .578, \eta_p^2 = 1.54 \times 10^{-3}$], suggesting that the speaker groups did not differ substantially in the extent to which they benefited from improved SNR conditions. Examination of effect size revealed a negligible interaction effect, implying that improved acoustic clarity produced relatively similar reductions in subjective listening effort across speaker groups.

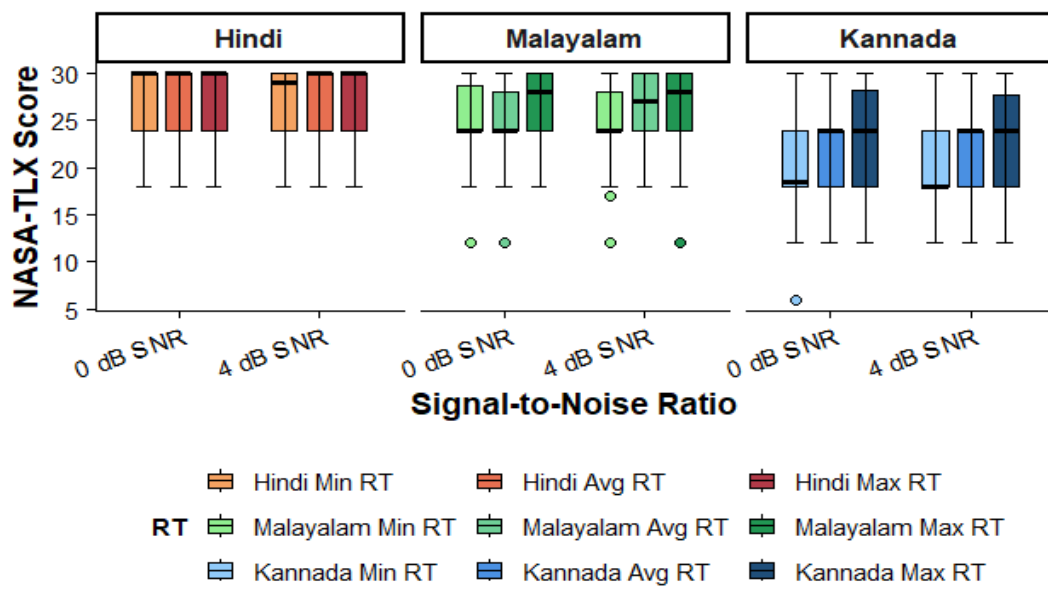
Importantly, the three-way interaction between RT, SNR, and Group was non-significant [$\chi^2(4) = 4.95, p = .292, \eta_p^2 = 2.30 \times 10^{-3}$]. This finding indicates that the combined effects of reverberation and noise on subjective listening effort did not vary substantially as a function of linguistic background. Examination of effect size revealed a negligible three-way interaction effect, suggesting that the combined influence of reverberation and noise produced relatively comparable patterns of subjective listening effort across native and non-native speaker groups.

4.4.2 Effect of Reverberation Time (RT) on NASA-TLX Scores

As LMER model revealed a differential and small effect of RT as well as its interaction with Group and SNR (Table 4.11), pairwise comparison with Bonferroni correction were carried out for comparison of NASA-TLX scores of three RT conditions for each group across SNRs. Figure 4.9 shows boxplots with median and IQR along with individual data points for NASA-TLX as a function of RT.

Figure 4.9

Boxplots illustrating NASA-TLX performance across SNR conditions for the Hindi, Malayalam, and Kannada speaker groups under Minimum, Average, and Maximum RT conditions. The central horizontal line within each box represents the median score, while the box boundaries indicate the interquartile range (IQR). Whiskers represent the spread of the data beyond the quartiles, and individual dots indicate each participant-performance and outliers.



Inspection of the figure 4.9 revealed clear RT-wise simple effects at each SNR $\times$ Group combination. Across both SNR conditions, subjective listening effort generally increased from Minimum RT to Maximum RT across all speaker groups. Hindi-speaking non-native Kannada speakers consistently demonstrated the highest NASA-TLX scores across all RT conditions, indicating greater perceived listening effort. In contrast, native Kannada speakers exhibited comparatively lower NASA-TLX scores across conditions, reflecting reduced cognitive workload during speech perception. Malayalam-speaking speakers showed intermediate performance between the Hindi and Kannada groups.

Under both 0 dB SNR and 4 dB SNR conditions, increases in reverberation time were associated with slight increases in NASA-TLX scores, particularly under Maximum RT conditions. This trend suggests that increased reverberation imposed greater perceptual and cognitive demands on speakers. The effect of RT appeared more prominent among the non-native speaker groups, especially the Hindi-speaking group, whereas native Kannada speakers demonstrated relatively lower listening effort even under adverse reverberation conditions.

The figure also demonstrates variability in NASA-TLX scores across participants. Kannada speakers showed comparatively wider score distributions and lower median values, whereas Hindi speakers exhibited consistently elevated scores with relatively smaller variability. Malayalam speakers demonstrated moderate variability with a few lower outlier values under some RT conditions.

Overall, the graphical pattern indicates that increasing reverberation elevated subjective listening effort across groups, while linguistic familiarity with Kannada reduced the cognitive demands associated with speech perception under adverse acoustic conditions. These findings visually support the significant main effect of RT and speaker group observed in the mixed-effects analysis.

Pairwise comparisons of RT conditions on NASA-TLX scores across SNR and speaker groups revealed differential effects of reverberation on subjective listening effort, as shown in Table 4.14. Overall, significant differences were more evident under the 4 dB SNR condition compared to the 0 dB SNR condition. Examination of effect sizes revealed predominantly negligible to moderate effects, indicating that reverberation produced

comparatively smaller influences on subjective listening effort than on objective task performance.

At 0 dB SNR for the Hindi group, none of the RT comparisons were statistically significant. The comparisons between Minimum RT and Average RT ($t = -0.55, p = 1.00, d = -0.08$), Minimum RT and Maximum RT ($t = -1.22, p = .66, d = -0.15$), and Average RT and Maximum RT ($t = -0.42, p = 1.00, d = -0.07$) demonstrated negligible to small effects, suggesting minimal changes in subjective listening effort across reverberation conditions.

For the Malayalam group at 0 dB SNR, significant differences were observed between Minimum RT and Maximum RT ($t = -5.25, p < .001, d = -0.64$) and between Average RT and Maximum RT ($t = -3.24, p = .004, d = -0.51$). Examination of effect sizes revealed moderate negative effects, implying that higher reverberation substantially increased perceived listening effort among Malayalam speakers. However, the comparison between Minimum RT and Average RT was non-significant ($t = -0.86, p = 1.00, d = -0.13$), indicating only a small effect.

Similarly, for the Kannada group at 0 dB SNR, the comparison between Minimum RT and Maximum RT was statistically significant ($t = -5.05, p < .001, d = -0.62$), demonstrating a moderate negative effect. The comparison between Average RT and Maximum RT demonstrated a small to moderate effect ($t = -2.15, p = .10, d = -0.34$), although it did not reach statistical significance following correction for multiple comparisons. The comparison between Minimum RT and Average RT remained non-significant ($t = -1.83, p = .20, d = -0.28$).

At 4 dB SNR, stronger RT-related effects were observed across groups. For the Hindi group, significant differences were observed between Minimum RT and Average RT ($t = -3.05, p = .007, d = -0.39$) and between Minimum RT and Maximum RT ($t = -3.44, p = .002, d = -0.38$), indicating small to moderate increases in listening effort with increasing

reverberation. However, the comparison between Average RT and Maximum RT was non-significant ($t = 0.10, p = 1.00, d = 0.01$).

For the Malayalam group at 4 dB SNR, significant differences were observed between Minimum RT and Average RT ($t = -5.66, p < .001, d = -0.73$) and between Minimum RT and Maximum RT ($t = -5.72, p < .001, d = -0.63$), demonstrating moderate negative effects. These findings suggest that reverberation substantially increased subjective listening effort among Malayalam speakers even under improved SNR conditions. In contrast, the comparison between Average RT and Maximum RT remained non-significant ($t = 0.74, p = 1.00, d = 0.10$).

Likewise, for the Kannada group at 4 dB SNR, significant differences were observed between Minimum RT and Average RT ($t = -3.98, p < .001, d = -0.51$) and between Minimum RT and Maximum RT ($t = -7.47, p < .001, d = -0.83$). Examination of effect sizes revealed moderate effects, with the largest effect observed for the Minimum RT versus Maximum RT comparison. The comparison between Average RT and Maximum RT did not reach statistical significance ($t = -2.34, p = .06, d = -0.31$).

Comparison of effect sizes across conditions revealed that the largest RT-related effects were observed in the Malayalam and Kannada groups under 4 dB SNR conditions, particularly for Minimum RT versus Maximum RT comparisons. In contrast, the Hindi group demonstrated comparatively smaller effects across RT conditions. Overall, Maximum RT conditions consistently produced higher NASA-TLX scores compared to Minimum RT conditions, indicating increased perceived listening effort with increasing reverberation.

Table 4.14*Pairwise Comparison of RT Conditions on NASA-TLX Performance*

SNR	RT	Hindi					Malayalam					Kannada				
		Estimate	SE	<i>t</i>	<i>p</i>	<i>d</i>	Estimate	SE	<i>t</i>	<i>p</i>	<i>d</i>	Estimate	SE	<i>t</i>	<i>p</i>	<i>d</i>
0 dB	Min vs Avg	-0.26	0.47	-0.55	1.00	N/A	-0.40	0.47	-0.86	1.00	N/A	-0.85	0.47	-1.83	0.20	N/A
	Min vs Max	-0.46	0.37	-1.22	0.66	N/A	-1.96	0.37	-5.25	< .001	-0.64	-1.88	0.37	-5.05	< .001	-0.62
	Avg vs Max	-0.20	0.48	-0.42	1.00	N/A	-1.56	0.48	-3.24	0.004	-0.51	-1.03	0.48	-2.15	0.10	N/A
4 dB	Min vs Avg	-1.20	0.39	-3.05	0.007	-0.39	-2.23	0.39	-5.66	< .001	-0.73	-1.57	0.39	-3.98	< .001	-0.51
	Min vs Max	-1.16	0.34	-3.44	0.002	-0.38	-1.93	0.34	-5.72	< .001	-0.63	-2.52	0.34	-7.47	< .001	-0.83
	Avg vs Max	0.04	0.41	0.10	1.00	N/A	0.30	0.41	0.74	1.00	N/A	-0.95	0.41	-2.34	0.06	N/A

Note: SNR -signal to noise ratio ; RT-reverberation time ; Min -minimum ; Avg-average ; Max -maximum

Interaction Effects Involving RT

The $RT \times SNR$ interaction was non-significant [$\chi^2(2) = 3.19, p = .203, \eta_p^2 = 5.10 \times 10^{-3}$], indicating that the effect of reverberation on NASA-TLX scores did not vary substantially across SNR conditions. Examination of effect size revealed a negligible interaction effect, suggesting that changes in reverberation produced relatively similar subjective listening effort ratings across both SNR conditions.

The $RT \times Group$ interaction was statistically significant [$\chi^2(4) = 11.18, p = .025, \eta_p^2 = 8.17 \times 10^{-3}$], suggesting that reverberation affected speaker groups differently. Examination of effect size revealed a negligible to small interaction effect, implying that increases in reverberation produced comparatively greater increases in subjective listening effort among non-native speakers relative to native Kannada speakers. Although all groups demonstrated changes in NASA-TLX scores across reverberation conditions, Malayalam and Hindi speakers tended to report relatively higher listening effort under adverse reverberation conditions.

Further, the three-way interaction between RT, SNR, and Group was non-significant [$\chi^2(4) = 4.95, p = .292, \eta_p^2 = 2.30 \times 10^{-3}$], indicating that the combined effects of reverberation and noise on subjective listening effort did not differ substantially across speaker groups. Examination of effect size revealed a negligible interaction effect, suggesting that the combined influence of RT and SNR produced relatively similar patterns of subjective listening effort across native and non-native speaker groups.

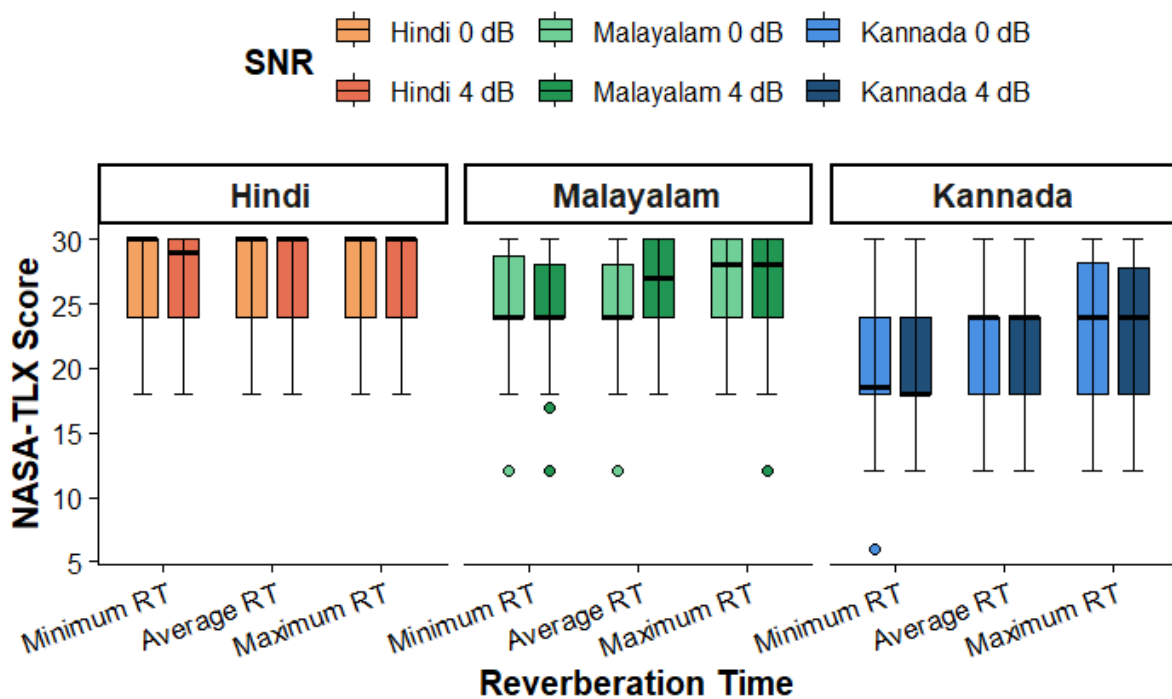
4.4.3 Effect of Signal-to-Noise Ratio (SNR) on NASA-TLX Scores

As LMER model revealed small but significant main effect of SNR, as well as its significant interaction with Group and RT (Table 4.11), pairwise comparison with Bonferroni correction were carried out for comparison of NASA-TLX scores of two SNR

conditions for each group across RTs. Figure 4.10 shows boxplots with median and IQR along with individual data points for NASA-TLX as a function of SNR.

Figure 4.10

Boxplots illustrating NASA-TLX performance across RT conditions (Minimum, Average, and Maximum RT) for Hindi, Malayalam, and Kannada speaker groups under 0 dB and 4 dB SNR conditions. The central horizontal line within each box represents the median score, while the box boundaries indicate the interquartile range (IQR). Whiskers represent the spread of the data beyond the quartiles, and individual dots indicate each participant-performance and outliers.



Inspection of Figure 4.10 revealed clear SNR-wise simple effects at each RT $\times$ Group combination. Overall, higher NASA-TLX scores were observed under the 0 dB SNR condition compared to the 4 dB SNR condition, indicating greater perceived listening effort in poorer signal-to-noise conditions.

Across all RT conditions, Hindi-speaking non-native Kannada speakers demonstrated consistently higher NASA-TLX scores than Malayalam-speaking and native Kannada speakers under both SNR conditions. Native Kannada speakers showed comparatively lower subjective listening effort across conditions. Malayalam-speaking speakers exhibited intermediate performance between the Hindi and Kannada groups.

The influence of SNR appeared more prominent under Average RT and Maximum RT conditions, where listening effort tended to increase further in the 0 dB SNR condition. In contrast, relatively lower NASA-TLX scores were observed in the 4 dB SNR condition, particularly under Minimum RT. These findings suggest that poorer signal clarity combined with increased reverberation elevated subjective listening effort, especially among non-native speakers.

Pairwise comparisons of SNR conditions on NASA-TLX scores across RT and speaker groups revealed relatively small effects of SNR on subjective listening effort, as shown in Table 4.15. Significant differences between 0 dB and 4 dB SNR conditions were observed only under selected RT and speaker group combinations. Examination of effect sizes revealed predominantly negligible to small effects, indicating that changes in SNR produced comparatively limited influences on perceived listening effort.

Under Minimum RT conditions, the Hindi group demonstrated a statistically significant difference between 0 dB and 4 dB SNR ($t = 2.36, p = .019, d = 0.26$), indicating a small effect. Similarly, the Kannada group also demonstrated a significant difference ($t = 2.21, p = .027, d = 0.24$), suggesting slightly lower NASA-TLX scores under the improved SNR condition. However, the Malayalam group did not demonstrate a significant difference ($t = 1.01, p = .315, d = 0.11$), indicating only a negligible effect of SNR.

Under Average RT conditions, the Malayalam group demonstrated a statistically significant difference between 0 dB and 4 dB SNR ($t = -2.93, p = .003, d = -0.49$), revealing a moderate negative effect. This finding suggests that improvements in SNR substantially reduced subjective listening effort among Malayalam speakers under moderate reverberation conditions. In contrast, the Hindi ($t = -0.30, p = .768, d = -0.05$) and Kannada ($t = 0.06, p = .956, d = 0.01$) groups demonstrated negligible and non-significant effects.

Under Maximum RT conditions, none of the speaker groups demonstrated statistically significant differences between SNR conditions. The Hindi group showed a negligible effect ($t = 0.24, p = .810, d = 0.03$), the Malayalam group demonstrated a small effect ($t = 1.00, p = .319, d = 0.12$), and the Kannada group also showed a negligible effect ($t = 0.30, p = .768, d = 0.04$). These findings indicate that improvements in SNR under highly reverberant conditions did not substantially reduce subjective listening effort.

Comparison of effect sizes across conditions revealed that the largest SNR-related effect was observed for the Malayalam group under Average RT conditions, followed by the Hindi and Kannada groups under Minimum RT conditions. Overall, the observed effect sizes ranged from negligible to moderate, suggesting that SNR exerted comparatively smaller effects on subjective listening effort relative to speaker group differences. These findings imply that linguistic familiarity contributed more strongly to perceived listening effort than improvements in acoustic clarity alone.

Table 4.15

Pairwise Comparison of SNR Conditions (0 dB vs 4 dB) on NASA-TLX Performance

RT	SNR	Hindi					Malayalam					Kannada				
		Estimate	SE	<i>t</i>	<i>p</i>	<i>d</i>	Estimate	SE	<i>t</i>	<i>p</i>	<i>d</i>	Estimate	SE	<i>t</i>	<i>p</i>	<i>d</i>
Min	0 dB vs 4 dB	0.79	0.34	2.36	0.019	0.26	0.34	0.34	1.01	0.315	N/A	0.74	0.34	2.21	0.027	0.24
Avg	0 dB vs 4 dB	-0.15	0.51	-0.30	0.768	N/A	-1.49	0.51	-2.93	0.003	-0.49	0.03	0.51	0.06	0.956	N/A
Max	0 dB vs 4 dB	0.09	0.37	0.24	0.810	N/A	0.37	0.37	1.00	0.319	N/A	0.11	0.37	0.30	0.768	N/A

Note: SNR -signal to noise ratio ; RT-reverberation time ; Min -minimum ; Avg-average ; Max -maximum

Interaction Effects Involving SNR

The $RT \times SNR$ interaction was non-significant [$\chi^2(2) = 3.19, p = .203, \eta_p^2 = 5.10 \times 10^{-3}$], indicating that the effect of SNR on NASA-TLX scores did not vary substantially across reverberation conditions. Examination of effect size revealed a negligible interaction effect, suggesting that improvements in acoustic clarity produced relatively similar changes in subjective listening effort across different RT conditions. Similarly, the $SNR \times Group$ interaction was non-significant [$\chi^2(2) = 1.10, p = .58, \eta_p^2 = 1.54 \times 10^{-3}$], indicating that the speaker groups did not differ substantially in the extent to which they benefited from improved SNR conditions.

Further, the three-way interaction between RT, SNR, and Group was non-significant [$\chi^2(4) = 4.95, p = .29, \eta_p^2 = 2.30 \times 10^{-3}$]. This finding indicates that the combined effects of reverberation and noise on subjective listening effort did not vary substantially across speaker groups. Examination of effect size revealed a negligible interaction effect, suggesting that the combined influence of RT and SNR produced relatively comparable patterns of perceived listening effort across all speaker groups.

4.5 Correlation between LEAP-Q Scores and Listening Effort Measures

Spearman correlation analysis was carried out to examine the relationship between second language proficiency, as measured using the LEAP-Q, and listening effort measures obtained from the dual-task paradigm i.e. primary and secondary task scores. In addition, spearman correlation was also measured between LEAP-Q scores and subjective workload ratings obtained using NASA-TLX. These were obtained separately for two SNRs (0 and 4 dB), and three RTs conditions (minimum, average, and maximum).

4.5.1. Correlation between LEAP-Q Scores with Primary and Secondary Task Performance

Figure 4.11 presents a heat map of Spearman's correlation coefficients, along with the corresponding p -values for significant relationships between LEAP-Q scores and performance on primary and secondary tasks across the panels. As shown in Figure 4.11, strong positive correlations were observed between LEAP-Q scores and both primary and secondary task performance across most RT and SNR conditions.

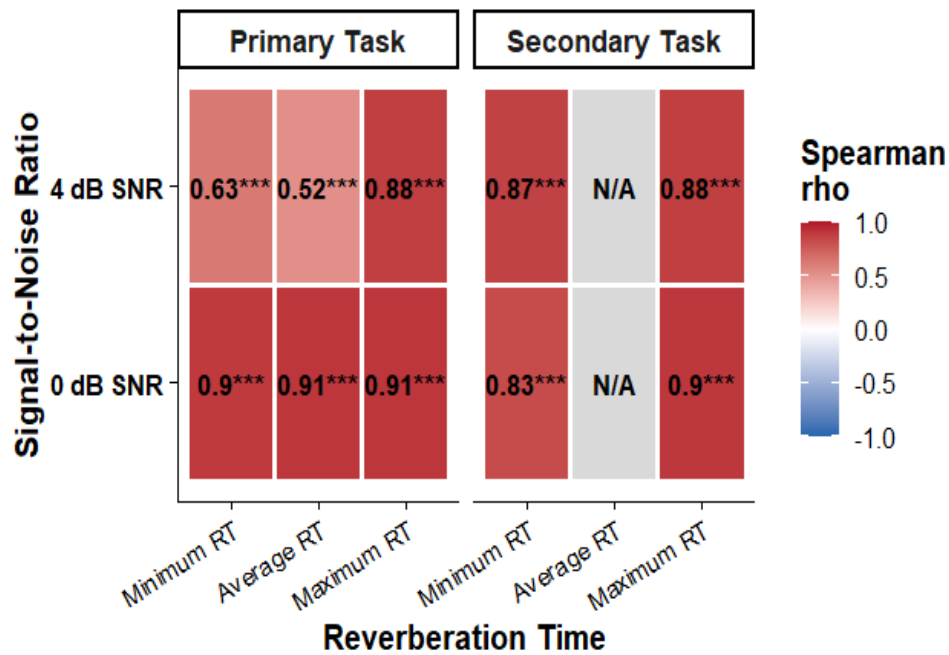
For the primary task, at 0 dB SNR, very strong positive correlations were observed across minimum RT ($r_s = 0.90, p < .001$), average RT ($r_s = 0.91, p < .001$), and maximum RT ($r_s = 0.91, p < .001$) with LEAP-Q scores. Similarly, at 4 dB SNR, significant positive correlations were observed at Minimum RT ($r_s = 0.63, p < .001$), Average RT ($r_s = 0.52, p < .001$), and Maximum RT ($r_s = 0.88, p < .001$).

For the secondary task, strong positive correlations were also observed between LEAP-Q scores and task performance. At 0 dB SNR, strong correlations were obtained at Minimum RT ($r_s = 0.83, p < .001$) and Maximum RT ($r_s = 0.90, p < .001$). At 4 dB SNR, similarly strong positive correlations were observed at Minimum RT ($r_s = 0.87, p < .001$) and Maximum RT ($r_s = 0.88, p < .001$). Correlation values for the Average RT condition were not computed due to methodological constraints and insufficient data. Specifically, blocks were categorized only as 'easy' or 'difficult,' where difficult blocks contained three average or maximum RT trials, while easy blocks consisted of three minimum RT trials (See method, section 3.4). This structure precluded meaningful calculation of average RT.

Overall, the heat map in Figure 4.11 demonstrates that higher Kannada language proficiency was consistently associated with better performance in both the primary and secondary tasks. The strongest correlations were observed under acoustically adverse conditions, particularly at 0 dB SNR and Maximum RT conditions.

Figure 4.11

Heat map showing Spearman correlation coefficients between LEAP-Q scores and primary and secondary task performance across different SNR and RT conditions.



Note: *** - $p < 0.001$

4.5.2 Correlation between LEAP-Q Scores and NASA-TLX Ratings

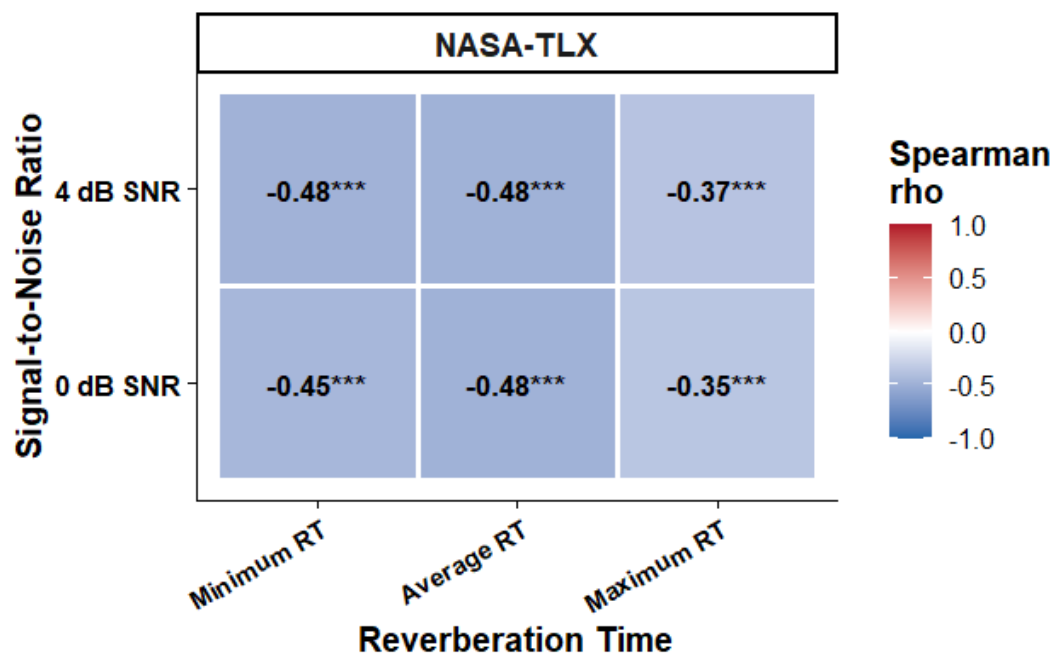
The relationship between LEAP-Q scores and subjective listening effort ratings obtained using NASA-TLX is presented in Figure 4.12. Moderate negative correlations were observed across all RT and SNR conditions.

At 0 dB SNR, significant negative correlations were observed at Minimum RT ($r_s = -0.45, p < .001$), Average RT ($r_s = -0.48, p < .001$), and Maximum RT ($r_s = -0.35, p < .001$).

Similarly, at 4 dB SNR, significant negative correlations were observed between LEAP-Q and NASA-TLX scores at Minimum RT ($r_s = -0.48, p < .001$), Average RT ($r_s = -0.48, p < .001$), and Maximum RT ($r_s = -0.37, p < .001$).

Figure 4.12

Heat map showing Spearman correlation coefficients between LEAP-Q scores and NASA-TLX ratings across different SNR and RT conditions.



Note: *** - $p < 0.001$

4.6 Results of test-retest reliability

To ensure the consistency of the measurements after one-month, 9 participants (10 % data) were re-tested on dual task paradigm measures and subjective ratings. Test-retest analysis was performed using Cronbach's alpha on Primary task, secondary task and NASA-TLX. The results are presented in Table 4.16, Table 4.17 and Table 4.18.

Table 4.16

Test–retest reliability Cronbach’s alpha values for all groups across SNR and RT conditions for primary task performance

SNR	RT	Hindi			Malayalam			Kannada		
		Test	Retest		Test	Retest		Test	Retest	
		Mean $\pm$ SD	Mean $\pm$ SD	α	Mean $\pm$ SD	Mean $\pm$ SD	α	Mean $\pm$ SD	Mean $\pm$ SD	α
0 dB	Min	20.33 $\pm$ 1.15	19.38 $\pm$ 2.08	0.83	22.33 $\pm$ 0.58	22.23 $\pm$ 0.58	0.98	24.00 $\pm$ 0.00	24.00 $\pm$ 0.00	0.99
0 dB	Avg	9.00 $\pm$ 1.00	9.93 $\pm$ 1.53	0.89	13.00 $\pm$ 1.00	12.24 $\pm$ 1.00	0.84	15.33 $\pm$ 0.58	15.00 $\pm$ 0.00	0.99
0 dB	Max	7.67 $\pm$ 1.53	7.33 $\pm$ 1.53	0.96	14.00 $\pm$ 1.00	13.67 $\pm$ 1.15	0.98	18.33 $\pm$ 1.15	18.33 $\pm$ 0.58	0.99
4 dB	Min	24.67 $\pm$ 0.58	24.00 $\pm$ 1.00	0.93	24.67 $\pm$ 0.58	24.33 $\pm$ 0.58	0.97	25.00 $\pm$ 0.00	25.00 $\pm$ 0.00	0.99
4 dB	Avg	16.67 $\pm$ 0.58	16.33 $\pm$ 0.58	0.98	16.33 $\pm$ 0.58	16.67 $\pm$ 0.58	0.98	17.00 $\pm$ 0.00	17.00 $\pm$ 0.00	1.00
4 dB	Max	18.53 $\pm$ 0.58	18.13 $\pm$ 1.00	0.97	19.67 $\pm$ 1.15	19.00 $\pm$ 0.00	0.97	21.67 $\pm$ 0.58	20.00 $\pm$ 0.00	0.98

Note. SNR = signal-to-noise ratio; RT = reverberation time; Min = minimum; Avg = average; Max = maximum; α = Cronbach’s alpha.

Table 4.17

Test–retest reliability Cronbach's alpha values for all groups across SNR and RT conditions for secondary task performance

SNR	RT	Hindi			Malayalam			Kannada		
		Test	Retest		Test	Retest		Test	Retest	
		Mean $\pm$ SD	Mean $\pm$ SD	α	Mean $\pm$ SD	Mean $\pm$ SD	α	Mean $\pm$ SD	Mean $\pm$ SD	α
0 dB	Min	11.00 $\pm$ 1.00	10.03 $\pm$ 0.58	0.88	12.33 $\pm$ 1.15	12.00 $\pm$ 1.73	0.98	15.00 $\pm$ 0.00	15.12 $\pm$ 0.00	0.96
0 dB	Max	4.00 $\pm$ 0.00	4.67 $\pm$ 0.58	0.96	7.00 $\pm$ 0.00	6.67 $\pm$ 1.53	0.98	8.67 $\pm$ 1.53	8.77 $\pm$ 1.15	0.98
4 dB	Min	16.33 $\pm$ 2.08	15.00 $\pm$ 1.73	0.89	18.00 $\pm$ 0.00	17.33 $\pm$ 0.58	0.97	19.33 $\pm$ 1.15	19.00 $\pm$ 1.00	0.97
4 dB	Max	10.00 $\pm$ 1.00	9.10 $\pm$ 1.00	0.87	10.67 $\pm$ 1.15	10.00 $\pm$ 0.00	0.97	14.00 $\pm$ 0.00	13.67 $\pm$ 0.58	0.98

Note. SNR = signal-to-noise ratio; RT = reverberation time; Min = minimum; Avg = average; Max = maximum; α = Cronbach's alpha.

Table 4.18

Test–retest reliability Cronbach’s alpha values for all groups across SNR and RT conditions for NASA–TLX scores

SNR	RT	Hindi			Malayalam			Kannada		
		Test	Retest		Test	Retest		Test	Retest	
		Mean $\pm$ SD	Mean $\pm$ SD	α	Mean $\pm$ SD	Mean $\pm$ SD	α	Mean $\pm$ SD	Mean $\pm$ SD	α
0 dB	Min	26.60 $\pm$ 3.91	24.67 $\pm$ 3.81	0.83	27.33 $\pm$ 0.90	26.00 $\pm$ 2.24	0.88	20.80 $\pm$ 3.10	22.33 $\pm$ 3.20	0.84
0 dB	Avg	27.50 $\pm$ 3.99	25.67 $\pm$ 4.03	0.86	27.00 $\pm$ 0.00	26.00 $\pm$ 2.37	0.86	22.00 $\pm$ 3.10	22.33 $\pm$ 3.39	0.89
0 dB	Max	26.67 $\pm$ 3.73	25.67 $\pm$ 3.85	0.88	27.17 $\pm$ 0.72	26.00 $\pm$ 2.26	0.81	22.00 $\pm$ 4.67	22.33 $\pm$ 3.23	0.77
4 dB	Min	25.83 $\pm$ 4.33	24.67 $\pm$ 3.79	0.76	26.83 $\pm$ 1.42	26.00 $\pm$ 2.22	0.84	20.00 $\pm$ 2.91	21.33 $\pm$ 3.18	0.86
4 dB	Avg	26.11 $\pm$ 4.54	24.67 $\pm$ 3.91	0.84	27.44 $\pm$ 0.53	26.00 $\pm$ 2.29	0.84	20.67 $\pm$ 3.16	22.33 $\pm$ 3.28	0.81
4 dB	Max	26.60 $\pm$ 3.98	24.67 $\pm$ 3.81	0.81	28.13 $\pm$ 0.52	27.00 $\pm$ 2.24	0.79	22.40 $\pm$ 4.79	22.33 $\pm$ 3.20	0.79

Note. SNR = signal-to-noise ratio; RT = reverberation time; Min = minimum; Avg = average; Max = maximum; α = Cronbach’s alpha.

The test–retest reliability analysis demonstrated good to excellent reliability across all three measures, namely the primary task, secondary task, and NASA–TLX scores, as indicated by the Cronbach’s alpha (α) values for all groups obtained across the different SNR and RT conditions. In general, Cronbach’s alpha values greater than 0.70 are considered acceptable, values above 0.80 indicate good reliability, and values above 0.90 indicate excellent reliability.

For the primary task performance (Table 4.16), the Cronbach’s alpha values ranged from 0.83 to 1.00 across the three groups and listening conditions, indicating good to excellent test–retest reliability. The Kannada group demonstrated the highest reliability overall, with alpha values consistently ranging from 0.98 to 1.00 across all SNR and RT conditions, suggesting highly stable performance between the test and retest sessions. The Malayalam group also showed strong reliability, with alpha values ranging from 0.84 to 0.98. Similarly, the Hindi group exhibited good reliability across conditions, with alpha values ranging from 0.83 to 0.98. These findings suggest that the primary task measure was highly consistent and reproducible across repeated administrations.

For the secondary task performance (Table 4.17), Cronbach’s alpha values ranged from 0.87 to 0.98, indicating good to excellent reliability across all groups and conditions. The Malayalam and Kannada groups demonstrated particularly strong reliability, with most alpha values above 0.97. The Hindi group showed slightly lower but still acceptable to excellent reliability values ranging from 0.87 to 0.96. The consistently high alpha values across the secondary task conditions indicate that the dual-task measure yielded stable and dependable responses across repeated testing sessions.

For the NASA–TLX scores (Table 4.18), the Cronbach’s alpha values ranged from 0.76 to 0.89, reflecting acceptable to good reliability across the groups and conditions. The Kannada group demonstrated relatively high reliability overall, with alpha values ranging

from 0.77 to 0.89. The Malayalam group showed alpha values between 0.79 and 0.88, while the Hindi group demonstrated alpha values ranging from 0.76 to 0.88. Although a few conditions showed comparatively lower alpha values, all values remained within the acceptable reliability range. These findings indicate that the NASA–TLX provided reasonably stable subjective workload ratings across repeated administrations.

Overall, the obtained Cronbach’s alpha values across all three measures support the strong test–retest reliability of the experimental procedures and outcome measures used in the present study. The findings suggest that the primary task, secondary task, and NASA–TLX measures were sufficiently stable and reliable for assessing listening effort across varying SNR and reverberation conditions in native and non-native Kannada speakers.

CHAPTER 5

DISCUSSION

The present study sought to investigate and compare the effects of signal-to-noise ratio (SNR), reverberation time (RT), and language background on listening effort in native Kannada speakers and two groups of non-native Kannada speakers (Hindi, & Malayalam as first language) using a dual-task paradigm and the NASA Task Load Index (NASA-TLX).

Prior to interpreting the main findings, the test-retest reliability of the measures was assessed in a subset of nine participants, three from each language group (Hindi, Malayalam, and Kannada) who completed all experimental conditions a second time after a one-month interval. Cronbach's alpha coefficients were computed for each measure across all RT and SNR conditions for each group. For the primary task, alpha values ranged from 0.83 to 1.0 across all conditions and groups, while in the secondary task, alpha values ranged from 0.87 to 0.98, both reflecting excellent reliability (Table 4.16, Table 4.17). For the NASA-TLX (subjective listening effort), alpha values ranged from 0.76 to 0.89, reflecting acceptable to good reliability across all conditions (Table 4.18). Across all three language groups, all RT conditions, and both SNR levels, the test and retest mean scores were closely comparable, with no systematic directional shift between sessions. Taken together, these results confirm that the primary task, secondary task, and NASA-TLX measures used in the present study demonstrate good to excellent stability over time, and that the observed group, RT, and SNR effects reported in the results chapter reflect genuine and reproducible differences in listening effort rather than measurement artefact.

The discussion that follows is organised around the seven major areas of findings: (1) the effect of language of speaker group on primary and secondary task performance, (2) the effect of reverberation time on primary and secondary task performance, (3) the effect of SNR on primary and secondary task performance, (4) the effect of group, RT, and SNR

on subjective listening effort (NASA-TLX), (5) interaction effects of RT, SNR, and Group on primary task, secondary task, and NASA-TLX, (6) the relationship between second language proficiency (LEAP-Q) and listening effort measures and (7) dissociation between objective and subjective measures of listening effort.

Each section discusses the observed findings, provides explanations for the findings, while drawing both supporting and contradicting evidences from the existing literature where relevant.

5.1 Effect of Language of the Speaker Group on Primary and Secondary Task Performance

One of the most consistent and robust findings of the present study was that native Kannada speakers outperformed both non-native groups (with Malayalam and Hindi as first language – L1) on both the primary (speech recognition) and secondary (recall/ memory) tasks across two SNR and three RT conditions. The Kannada group demonstrated the highest primary task scores (Table 4.3), reflecting superior speech recognition accuracy, while simultaneously exhibiting the highest secondary task scores (Table 4.8), which are indicative of lower cognitive listening effort in the dual-task framework. The Hindi group consistently showed the lowest performance on both measures, with the Malayalam group occupying an intermediate position throughout. (Table 4.3, Table 4.8 and Figure 4.1)

These group differences in primary task performance reflect well-documented phenomena in the speech perception literature concerning native and non-native speakers. It is widely established that non-native speakers experience greater difficulty understanding speech in adverse acoustic environments compared to native speakers, even when their proficiency in the target language is relatively high (Borghini & Hazan, 2018; Lecumberri et al., 2010; Meador et al., 2000; van Wijngaarden et al., 2002). Takata & Nábělek (1990) demonstrated in one of the earlier systematic investigations that non-native Japanese

speakers obtained significantly lower speech recognition scores in noise compared to their native American counterparts, despite performing comparably in quiet conditions. This native-speaker advantage under degraded conditions has been replicated extensively and forms a robust backdrop against which the current findings can be situated. In the present study, both Malayalam and Hindi speakers showed lower recognition accuracy, with the degradation becoming steeper as RT and noise conditions worsened, suggesting that non-native speakers are disproportionately disadvantaged in complex acoustic environments.

The mechanism underlying this disadvantage is generally attributed to the reduced automaticity and lexical efficiency of second language (L2) processing. When listening in an L2, speakers must rely more heavily on effortful, controlled processes rather than the fast, automatic speech decoding that characterises native-language listening (Lecumberri et al., 2010). In adverse conditions, where the acoustic signal is already compromised, this additional processing load can overwhelm available cognitive resources, leading to both poorer recognition accuracy and increased listening effort. Borghini & Hazan (2018) confirmed, using physiological measures of pupil dilation, that non-native English speakers demonstrated objectively higher listening effort compared to native speakers even when recognition accuracy was similar, underscoring the important distinction between performance and effort.

The gradient between groups with Hindi speakers showing the lowest performance relative to Malayalam speakers who fell in between native Kannada and Hindi groups is a particularly informative finding (Table 4.3, Table 4.8 and Figure 4.1). This gradient likely reflects the degree of linguistic proximity between the L1 and Kannada. Kannada and Malayalam are both Dravidian languages and share substantial phonological, morphological, and syntactic features including similar consonant inventory, vowel structure, and agglutinative morphology (Krishnamurti, 2003). Malayalam speakers

listening to Kannada would therefore encounter a phonological landscape that is considerably more familiar than what Hindi speakers face. Hindi, an Indo-Aryan language, differs markedly from Kannada in its segmental phonology, prosodic structure, and syntactic organisation. Malayalam exhibits a relatively denser and more contrast-rich phonological system compared to Hindi, particularly in its vowel inventory and segmental distinctions (Geethakumary, 2002). Malayalam possesses a relatively rich phonological system with phonemic vowel length and a larger segmental inventory compared to Hindi, which has a more symmetrical and relatively smaller vowel system ((Namboodiripad, 2016; Ohala & Ohala, 1992) According to dispersion theory, vowel systems are organized to maximize perceptual distinctiveness, and languages with larger inventories tend to exhibit greater contrast density in acoustic space (Liljencrants et al., 1972; Schwartz et al., 1997). Given the close phonological and typological similarity between Malayalam and Kannada as Dravidian languages, Malayalam speakers show better mapping of Kannada phonetic categories. The greater phonological and typological distance between Hindi and Kannada likely meant that Hindi speakers had a less efficient phonological parsing system when processing Kannada speech, resulting in greater cognitive effort and poorer speech recognition accuracy.

This explanation is consistent with the linguistic distance hypothesis, which proposes that the degree of structural similarity between an L1 and a target language influences both the ease of acquisition and the cognitive demands of real-time processing in that language (Ringbom, 2006). Shastri & Kumar (2015), working within an Indian multilingual context, found that recognition of Malayalam phone contrasts was influenced by the linguistic relationship between the speaker's L1 and the target language, with greater typological distance producing poorer phonological identification, in native speakers of Hindi and native speakers of Malayalam population. This reflects how shared phonological

features between two Dravidian languages (Malayalam and Kannada) create both advantages and interference in processing. While that study did not focus on listening effort per se, it supports the broader point that linguistic background shapes speech perception outcomes in meaningful ways. The current data extend this line of work by demonstrating that it is not just recognition accuracy but also the cognitive load associated with listening that is systematically influenced by the nativity and typological proximity of the speaker's L1 to the target language. Manikandan & Geetha (2020) working within an Indian multilingual context, found that Kannada sentence recognition scores were significantly poorer in the presence of native (Kannada) two-talker babble compared to non-native babble conditions (Tamil, Telugu, Hindi), in 20 younger native Kannada-English bilingual adults with normal hearing, demonstrating that linguistic similarity between the target and masker language plays a primary role in informational masking, consistent with the broader finding that linguistic background shapes speech perception outcomes.

The secondary task results reinforced this interpretation. In the dual-task paradigm, secondary task performance is understood as an inverse proxy for listening effort. When speakers must allocate more cognitive resources to the primary task, fewer resources are available for the secondary task, and secondary task performance deteriorates (Gosselin & Gagné, 2011). The Kannada group demonstrated consistently higher secondary task scores, indicating that they required fewer cognitive resources to achieve their (already superior) primary task performance. In contrast, the Hindi group showed the lowest secondary task scores, reflecting the greatest cognitive load. These effects were large and consistent across all acoustic conditions, suggesting that the native-language advantage operates as a relatively stable facilitator of efficient processing rather than a context-dependent benefit.

Peng & Wang (2019) observed a comparable pattern in a study of non-native Chinese speakers processing English, finding that non-native speakers exhibited higher dual-task

reaction times compared to native speakers in noise, and that this difference was attributable to the additional cognitive resources demanded by L2 processing. Similarly, Brännström et al. (2021) reported that non-native children experienced higher listening effort and fatigue than native peers in school listening conditions. These findings corroborate the view that native language status is a robust predictor of cognitive listening load, and the current results add to this by showing that the specific typological relationship between the L1 and the target language (Dravidian versus Indo-Aryan) further modulates the magnitude of these differences within a non-native group. This has practical significance for multilingual environments like India, where individuals may regularly need to communicate in languages that differ substantially from their mother tongue.

5.2 Effect of Reverberation Time on Primary and Secondary Task Performance

The present study found a significant main effect of RT on both primary and secondary task performance, with increasing reverberation time consistently degrading speech recognition accuracy and increasing cognitive listening effort across all speaker groups and SNR conditions (Table 4.4, Table 4.9) Primary task scores followed a clear gradient from Minimum RT to Maximum RT, with performance being highest under Minimum RT, intermediate under Average RT, and lowest under Maximum RT (Figure 4.3). Pairwise comparisons confirmed that differences between all three RT conditions were statistically significant across groups, with very large effect sizes (Table 4.4).

These findings align closely with the existing literature on the effects of reverberation on speech perception. Reverberation is known to degrade the temporal fine structure and spectral contrast of the speech signal by allowing earlier reflections to overlap with and smear subsequent sounds, a phenomenon known as temporal masking and spectral smearing (Nabelek & Pickett, 1974). This degradation is particularly harmful to consonant

recognition, where fine-grained temporal cues carry phonemic distinctions. As RT increases, the accumulation of overlapping reflections progressively blurs these cues, making it increasingly difficult for speakers to accurately identify target words. Finitzo-Hieber & Tillman (1978) conducted a foundational study with normal-hearing speakers and found that increasing RT from 0.0 s to 1.2 s significantly degraded speech recognition scores for monosyllabic English words, consistent with the findings of the present study where the Maximum RT condition (4.37 s) yielded significantly lower speech perception in Kannada sentences compared to Average RT (1.66 s) and Minimum RT (0.74 s) conditions in both native and non-native speakers.

The secondary task results further demonstrated that increasing RT elevated cognitive listening effort substantially. The beta coefficient for Maximum RT in the LMER model was very large ($\beta = -1.32$), indicating that reverberation exerted the strongest single influence on secondary task performance among all the predictors examined (Table 4.6). This is a noteworthy finding because it establishes that reverberation, even without the concurrent presence of background noise, places a significant cognitive burden on speakers. The effort required to compensate for the perceptual degradation caused by reverberation draws on attentional and working memory resources that would otherwise be available for secondary task performance. In contrast to this finding, Picou et al. (2016) using a semantic dual-task paradigm found that while noise increased cognitive load as measured by reaction time of secondary task, reverberation did not significantly increase listening effort even though it reduced speech intelligibility. This divergence with the current results warrants consideration. One possible explanation lies in the substantial differences in the degree of reverberation employed across the two studies. Picou et al. (2016) used three RT conditions i.e $RT_{30} < 100$ ms, 473 ms, and 875 ms — with the most reverberant condition reaching only 875 ms. In the present study, by contrast, even the Minimum RT condition was 740 ms, the

Average RT was 1,660 ms, and the Maximum RT was 4,370 ms — approximately five times longer than the most reverberant condition tested by Picou et al. (2016). These comparatively mild reverberation levels may not have imposed the sustained perceptual challenge required to produce measurable secondary task decrements in a dual-task paradigm. In the present study, the range of RT conditions was likely more ecologically representative of real Indian classroom and public environments, where reverberation times can be substantially elevated due to inadequate acoustic treatment (Kanigalpula et al., 2026; Saravanan et al., 2019). It is also possible that the use of a non-native target language in the current study increased sensitivity to reverberation, since non-native speakers have fewer compensatory resources to offset signal degradation.

5.3 Effect of Signal-to-Noise Ratio on Primary and Secondary Task Performance

The study demonstrated a highly significant effect of SNR on both primary and secondary task performance. Across all speaker groups and RT conditions, primary task scores were consistently higher at 4 dB SNR compared to 0 dB SNR, confirming that improved signal clarity facilitated better speech recognition. Similarly, secondary task performance was significantly better at 4 dB SNR, indicating reduced cognitive listening effort under more favourable noise conditions (Table 4.5, Table 4.10).

These findings are consistent with a well-established body of research demonstrating that increasing background noise degrades speech intelligibility and elevates listening effort in a dose-dependent manner (Brown & Strand, 2019; Strand et al., 2018). As SNR decreases, the acoustic signal becomes increasingly masked, forcing speakers to invest greater cognitive resources in extracting phonological and linguistic information from a degraded input. The dual-task paradigm captures this resource allocation process by showing that as primary task demands increase with worsening SNR, secondary task performance suffers. This pattern is clearly reflected in the present data.

A particularly interesting observation is that the magnitude of the SNR benefit differed across speaker groups. The pairwise comparisons revealed that the SNR-related improvement in secondary task performance was largest for the Hindi group (effect size $d = -2.70$ under Minimum RT) and smallest for the Kannada group ($d = -0.90$ under Minimum RT) (Table 4.10). This pattern indicates that non-native speakers, particularly those with the greatest typological distance from Kannada derived greater absolute benefit from improved SNR. This can be understood through the lens of resource competition: non-native speakers are already operating at or near the limits of their cognitive capacity at the poorer SNR. When SNR improves, they experience proportionally greater relief because the reduced perceptual demand frees up resources that were previously monopolised by the effort of decoding a degraded L2 signal. Native Kannada speakers, operating closer to their performance ceiling, had less room to benefit.

This interpretation resonates with findings from Kilman et al. (2015), who examined speech perception in noise across different masker types in non-native speakers and found that non-native speakers required more favourable SNRs to achieve the same intelligibility levels as native speakers. Alqattan & Turner (2021) similarly found that bilingual speakers required better SNRs compared to monolinguals on QuickSIN measures, suggesting a systematic disadvantage for non-native speakers under adverse noise conditions. The present study extends this finding by showing that the relationship between SNR and cognitive load not just accuracy, but it is also modulated by native language status.

5.4 Effect of Group, Reverberation, and SNR on Subjective Listening Effort (NASA-TLX)

The pattern of results for subjective listening effort as measured by the NASA-TLX diverged in several interesting ways from the objective dual-task findings, offering insights into the multidimensional nature of listening effort.

Consistent with the dual-task results, speaker group emerged as the strongest predictor of NASA-TLX scores ($\chi^2(2) = 56.76$, $p < .001$, $\eta_p^2 = 0.38$). Native Kannada speakers reported the lowest subjective listening effort, followed by Malayalam speakers, with Hindi speakers reporting the highest perceived workload (Table 4.13). The effect sizes were large, indicating that this gradient in perceived effort was not a subtle statistical artefact but a robust perceptual reality for the participants. This is consistent with the theoretical expectation that linguistic familiarity should reduce both objective and subjective effort, as the FUEL framework (Pichora-Fuller et al., 2016) emphasises that perceived effort is shaped by the interaction of sensory demands, cognitive capacity, and linguistic competence. For native Kannada speakers, automatic phonological and lexical processing of Kannada speech minimised the cognitive demands of the primary task, thereby reducing the mental demand, effort, and frustration subscales of the NASA-TLX that collectively reflect perceived listening workload. This explains the significantly lower NASA-TLX scores observed for the Kannada group across all RT and SNR conditions in the present study.

The subjective superiority of the Kannada group is further interpretable in terms of the ease with which they could form and access lexical representations of Kannada words. For native speakers, word recognition is largely automatic and requires little conscious monitoring, reducing the subjective sense of effort. Non-native speakers, especially Hindi speakers, are aware of the additional cognitive work they must engage in to successfully comprehend a second language, and this awareness is likely reflected in their higher NASA-TLX ratings. Francis et al. (2021) found that individual differences in language background influenced both subjective and physiological measures of listening effort, though not always

in the same direction, supporting the idea that subjective and objective indices capture related but partially distinct aspects of the effort experience.

The non-significant main effect of RT on NASA-TLX scores ($\chi^2(2) = 1.51, p = .47$) is one of the important findings of the study. Despite the fact that RT exerted the strongest influence on objective dual-task performance, participants did not perceive increasing reverberation as significantly increasing their workload. This dissociation between objective and subjective measures of effort has been noted in previous research. Picou et al. (2016) similarly reported that reverberation increased intelligibility difficulty without producing the same level of increase in perceived cognitive load as noise. One plausible explanation is that speakers may have limited metacognitive awareness of the cognitive demands specifically imposed by reverberation, as opposed to noise. Background noise is typically experienced as an obvious and salient interference, whereas reverberation produces a subtler form of signal degradation such as a blurring or lengthening of sounds, which may be less consciously attributable to increased mental workload. There is also the possibility of an anchoring or habituation effect. If participants experienced reverberation throughout testing as a relatively constant background condition (compared to the more acute contrast created by noise level changes), they may have calibrated their subjective ratings against an already-elevated reverberant baseline, reducing their sensitivity to RT-related changes. This mirrors findings by Lau et al. (2019) who demonstrated that physiological and subjective measures of effort do not always converge, particularly when the source of degradation is less perceptually salient.

Importantly, when examining pairwise RT comparisons on NASA-TLX, Minimum RT versus Maximum RT comparisons were significant for the Malayalam and Kannada groups, particularly under the 4 dB SNR condition (e.g., $r_s = -0.63$ to -0.83) (Table 4.14). This suggests that while the main effect of RT was not global, the extreme reverberant

condition (Maximum RT) did produce perceptible increases in effort for some groups, particularly when combined with better SNR. Under better SNR conditions, speakers may be more sensitised to residual signal distortions caused by reverberation, since noise is no longer the dominant challenge. This could explain why RT effects on NASA-TLX were more apparent under 4 dB SNR than under 0 dB SNR.

The significant but small effect of SNR on NASA-TLX ($\chi^2(1) = 5.55, p = .019$) confirmed that improving SNR modestly reduced perceived workload, but the effect size was negligible. This aligns with the finding by Strand et al. (2018) that noise substantially increases listening effort, though the SNR range examined in the present study (0 vs. 4 dB) was relatively narrow and may not have been sufficient to elicit large subjective differences. The effect was perhaps most apparent in individual comparisons: the Malayalam group showed a significant and moderate reduction in NASA-TLX under Average RT when SNR improved to 4 dB ($d = -0.49$), suggesting that for non-native speakers at moderate reverberation, improved signal clarity does provide subjective relief from cognitive workload.

5.5 Interaction Effects of RT, SNR, and Group on Primary Task, Secondary Task, and NASA-TLX

Beyond the main effects of RT, SNR, and speaker group, the present study yielded a series of important interaction effects that shed light on how acoustic and linguistic factors jointly shape listening effort. These interaction effects are discussed below for each of the three listening effort measures (primary task, secondary task, and NASA-TLX).

5.5.1 Interaction Effects on Primary Task Performance

The $RT \times SNR$ interaction on primary task performance was statistically significant [$\chi^2(2) = 142.99, p < .001, \eta_p^2 = 4.27 \times 10^{-8}$] (Table 4.2), indicating that the benefit of

improved SNR on speech recognition varied as a function of the level of reverberation. Specifically, the improvement in primary task scores when SNR increased from 0 dB to 4 dB was larger under Minimum RT conditions than under Maximum RT conditions. This interaction reflects a well-recognised limitation of SNR improvements in highly reverberant environments: when RTs are long, the temporal smearing of the speech signal reduces the effective gain obtained from even a favourable SNR, because the reverberant energy itself continues to degrade spectral and temporal cues independently of the direct-sound-to-noise relationship (Nabelek & Pickett, 1974). As a result, improving SNR is a more effective strategy for enhancing speech recognition in low-reverberation environments than in highly reverberant spaces. This finding is consistent with Rennie et al. (2014), who demonstrated that the combined degradation of noise and reverberation produces non-additive impairments in speech intelligibility, and that the presence of reverberation constrains the benefit that can be obtained from improvements in SNR alone.

The $RT \times \text{Group}$ interaction on primary task performance was highly significant [$\chi^2(4) = 492.69, p < .001, \eta_p^2 = 1.45 \times 10^{-8}$] (Table 4.2), indicating that the detrimental effects of increasing reverberation on speech recognition were not uniform across language groups. Although all three groups showed progressively lower primary task scores as RT increased, the non-native groups, particularly the Hindi group demonstrated greater reductions in performance under Average RT and Maximum RT compared to native Kannada speakers. The large effect size associated with this interaction implies that linguistic background moderated the impact of reverberation on speech recognition. This differential susceptibility is theoretically coherent within the framework of top-down processing compensation (Mattys et al., 2012) as reverberation degrades the acoustic signal, speakers who possess robust and automatic phonological representations of the target language can more effectively restore or predict distorted segments using stored lexical and phonological

knowledge. Native Kannada speakers, for whom lexical access is more automatic, are better positioned to exercise such top-down compensation. In contrast, non-native speakers, whose phonological representations of Kannada are less well-specified and less automatically accessed, cannot compensate as efficiently for reverberation-induced perceptual degradation, resulting in steeper performance decrements. Peng & Wang (2019) observed a similar pattern, finding that non-native speakers showed disproportionately greater listening difficulty than native speakers when both noise and reverberation were present simultaneously, attributing this to the compounded demands on cognitive resources when L2 processing inefficiency co-occurs with adverse acoustics.

The SNR $\times$ Group interaction on primary task performance was also highly significant [$\chi^2(2) = 389.58, p < .001, \eta_p^2 = 3.70 \times 10^{-8}$] (Table 4.2), demonstrating that the degree of benefit obtained from improved SNR varied across language groups. Native Kannada speakers showed proportionally smaller gains in primary task scores when SNR improved from 0 dB to 4 dB compared to the non-native groups, particularly under conditions of minimum reverberation. This finding reflects a performance ceiling effect for native speakers, who were already performing near the upper limit of their recognition capacity under 0 dB SNR at Minimum RT, leaving less room for improvement at 4 dB SNR. In contrast, non-native speakers, operating well below ceiling even under Minimum RT conditions, had substantially more capacity to benefit from SNR improvements. Kilman et al. (2014) reported a comparable finding that non-native speakers derived greater absolute improvements in speech recognition scores from favourable SNR conditions compared to native speakers, whose performance showed smaller SNR-related gains due to near-ceiling performance. These differential interaction effects have important practical implications in acoustically challenging settings: acoustic interventions that improve SNR are likely to yield

disproportionate benefits for non-native speakers, who stand to gain the most from even modest improvements in signal clarity.

The three-way interaction between RT, SNR, and Group on primary task performance was statistically significant [$\chi^2(4) = 27.24, p < .001, \eta_p^2 = -9.35 \times 10^{-5}$] (Table 4.2), indicating that the combined effects of reverberation and noise on speech recognition differed meaningfully across language groups. Specifically, the three-way interaction terms for Average RT $\times$ 4 dB SNR $\times$ Malayalam ($\beta = -0.49, p = .005$) and Maximum RT $\times$ 4 dB SNR $\times$ Malayalam ($\beta = -0.56, p = .002$) revealed moderate negative effects, indicating that Malayalam speakers experienced additional reductions in recognition accuracy under the combined influence of reverberation and noise beyond what would be predicted from the main effects alone. In contrast, the corresponding interaction terms for the Kannada group were non-significant, confirming that native Kannada speakers were comparatively less vulnerable to the compound degradation imposed by co-occurring noise and reverberation. This three-way interaction suggests that the adverse effects of noise and reverberation are not simply additive but rather interact in a manner that is further modulated by the linguistic competence of the speaker. The finding is theoretically interpretable within a shared resource competition model (Pelle, 2018), in which the simultaneous demands of reverberation-induced temporal smearing, noise-induced spectral masking, and L2 processing inefficiency place greater cumulative strain on limited cognitive resources for non-native speakers than for native speakers. It also aligns with the position advanced by Mattys et al. (2012) that adverse conditions compound in their effects on speech recognition, with the compounding being most severe for speakers who already have fewer top-down compensatory resources available as experienced by non-native speakers in the present study.

5.5.2 Interaction Effects on Secondary Task Performance

The $RT \times SNR$ interaction on secondary task performance was not statistically significant [$\chi^2(1) = 0.39, p = .53$] (Table 4.7). The non-significant finding indicates that the effects of reverberation and SNR on cognitive listening effort were largely independent of one another on the secondary task. This is a somewhat unexpected result given the significant $RT \times SNR$ interaction on primary task performance (Table 4.2) and suggests that while the two acoustic factors interact in their effects on speech recognition accuracy, their combined influence on the downstream allocation of cognitive resources to the secondary task may operate through more additive pathways. Picou et al. (2016) similarly found that in their dual-task paradigm, the interaction between noise and reverberation on secondary task response times was not significant, suggesting that these two factors may independently and separately affect different stages of speech processing and cognitive resource allocation. However, this interpretation must be treated with caution; a more powerful design with a wider range of acoustic conditions might reveal a significant $RT \times SNR$ interaction on secondary task performance.

The $RT \times Group$ interaction on secondary task performance was statistically significant [$\chi^2(2) = 7.53, p = .023, \eta_p^2 = 0.003$] (Table 4.7), indicating that the effect of reverberation on cognitive listening effort differed across speaker groups, though the effect size was very small. While all three groups showed increased cognitive load with increasing RT, non-native speaker groups demonstrated comparatively greater increases in secondary task effort under Maximum RT conditions. At 0 dB SNR, the effect size for the Minimum RT versus Maximum RT comparison on the secondary task was largest for the Hindi group ($d = 4.80$) and smallest for the Kannada group ($d = 4.10$), with the Malayalam group falling in between ($d = 4.60$). Although the absolute differences between groups were modest, this graded pattern of RT susceptibility aligns with the expectation derived from the FUEL framework (Pichora-Fuller et al., 2016) that speakers with greater sensory and linguistic

processing challenges will show greater cognitive load under adverse conditions. Non-native speakers processing Kannada in reverberation must simultaneously manage the perceptual degradation caused by temporal smearing and the inherent cognitive inefficiency of L2 phonological processing, placing a greater cumulative strain on working memory and attentional resources. This interpretation is supported by Brännström et al. (2021) who found that non-native children showed greater listening effort than native peers in challenging classroom acoustic conditions, with differences becoming more pronounced as listening difficulty increased.

The SNR $\times$ Group interaction on secondary task performance was also statistically significant [$\chi^2(2) = 8.44, p = .02, \eta_p^2 = 0.02$] (Table 4.7), indicating that the degree to which speakers benefited from improved SNR differed across language groups. As with the primary task, non-native speakers, particularly the Hindi group derived greater absolute cognitive relief from improved SNR, as reflected in larger secondary task improvements when SNR increased from 0 dB to 4 dB. Under Minimum RT, the Hindi group demonstrated a secondary task effect size of $d = -2.70$ for the 0 dB vs. 4 dB SNR comparison, compared to $d = -0.90$ for the Kannada group. This asymmetry in SNR-related benefit reflects the greater cognitive vulnerability of non-native speakers at adverse SNRs. Because Hindi speakers are already expending substantially more cognitive resources at 0 dB SNR to process L2 Kannada speech, even a modest 4 dB improvement in SNR provides proportionally greater relief from cognitive load compared to native Kannada speakers, who are operating more efficiently across the full range of SNR conditions. This finding is consistent with the theoretical position advanced by Rönnberg et al. (2013) in the Ease of Language Understanding (ELU) model, which predicts that speakers with greater phonological processing inefficiency will show more pronounced cognitive responses to

changes in signal quality, since their processing relies more heavily on effortful, explicit mechanisms that are particularly sensitive to variations in acoustic clarity.

The three-way interaction between RT, SNR, and Group on secondary task performance was statistically significant and yielded a moderate effect size [$\chi^2(2) = 19.38$, $p < .001$, $\eta_p^2 = 0.07$] (Table 4.7), indicating that the combined effects of reverberation and noise on cognitive listening effort differed meaningfully across speaker groups. The three-way interaction coefficient for Maximum RT $\times$ 4 dB SNR $\times$ Malayalam ($\beta = -0.43$, $p < .001$) and Maximum RT $\times$ 4 dB SNR $\times$ Kannada ($\beta = -0.29$, $p = .004$) revealed moderate negative effects, indicating that under the combination of maximum reverberation and improved SNR, these groups showed additional reductions in secondary task performance beyond what would be predicted from the individual main effects and two-way interactions. This finding implies that when reverberant conditions are at their most extreme, even improvements in SNR do not fully compensate for the cognitive burden imposed by reverberation, and that this compound challenge is experienced most acutely by non-native speakers. The specific pattern of the three-way interaction, wherein Malayalam and Kannada groups, but not the Hindi group, show significant interaction effects may seem counterintuitive at first. However, it is interpretable in terms of a differential processing advantage and baseline: at Maximum RT and 4 dB SNR, Malaysian and Kannada speakers benefit more from SNR improvement under lower RT but are disproportionately challenged when Maximum RT co-occurs with any SNR condition, suggesting these groups are more sensitive to the acoustic contrast between RT conditions. Hindi speakers, already at a high cognitive load baseline across conditions, show comparatively less differentiation across RT $\times$ SNR combinations. This interpretation parallels findings from Peng & Wang (2019), who demonstrated that the compounding of noise and reverberation produced non-linear

increases in listening effort for non-native speakers, with the most adverse compound conditions revealing the largest group differences in cognitive load.

5.5.3 Interaction Effects on NASA-TLX (Subjective Listening Effort)

The $RT \times SNR$ interaction on NASA-TLX scores was not statistically significant [$\chi^2(2) = 3.19, p = .20, \eta_p^2 = 5.10 \times 10^{-3}$] (Table 4.12), indicating that the subjective experience of listening effort did not vary substantially as a joint function of reverberation and SNR. This non-significant finding mirrors the pattern observed for the $RT \times SNR$ interaction on the secondary task and suggests that the interaction between these two acoustic variables on cognitive load, while present on objective performance measures, is not perceptibly experienced by speakers in a manner that they can consciously report via self-rating scales. One reason for this may be that the NASA-TLX is a retrospective, holistic measure of workload rather than a real-time index of cognitive resource allocation, and speakers may integrate their experience of noise and reverberation into a single global sense of effort rather than perceiving their combined effects as multiplicative or interactive. This interpretation is consistent with Francis & Love (2020) who argued that subjective effort measures are primarily shaped by the affective and motivational dimensions of the listening experience, which may be less sensitive to the specific acoustic interaction patterns captured by behavioural performance measures.

The $RT \times Group$ interaction on NASA-TLX was statistically significant [$\chi^2(4) = 11.18, p = .03, \eta_p^2 = 8.17 \times 10^{-3}$] (Table 4.12), though with a negligible to small effect size, indicating that while the effect of reverberation on perceived effort differed across speaker groups, the magnitude of this difference was modest. The interaction coefficients for Maximum $RT \times Malayalam$ ($\beta = 0.31, p = .004$) and Maximum $RT \times Kannada$ ($\beta = 0.30, p = .007$) indicated that these groups showed greater increases in perceived effort under

Maximum RT compared to the Hindi group. This unexpected pattern where the groups with lower baseline effort (Malayalam and Kannada) showed more pronounced RT-related increases in NASA-TLX than the highest-effort group (Hindi) is interpretable in terms of a ceiling or anchoring effect. Hindi group who reported consistently high NASA-TLX scores across all conditions, may have reached a perceived effort ceiling that limited their ability to register further increases in effort as RT worsened. Malayalam and Kannada group beginning from lower baseline effort levels, had more perceptual ‘headroom’ to detect and report increases in workload as reverberation increased to its maximum level. A comparable ceiling effect on subjective effort ratings was reported by Lau et al. (2019), who observed that physiological and subjective measures diverged most strongly in participants whose subjective ratings were already at or near the upper end of the rating scale, with subjective measures showing less sensitivity to condition changes at high-effort baselines.

The SNR $\times$ Group interaction on NASA-TLX was not statistically significant [$\chi^2(2) = 1.10, p = .58, \eta_p^2 = 1.54 \times 10^{-3}$] (Table 4.12), indicating that the speaker groups did not differ substantially in the extent to which improved SNR reduced their perceived listening effort. This null finding contrasts with the significant SNR $\times$ Group interactions observed for both primary and secondary task performance, and suggests that while linguistic background modulates the objective cognitive benefit of improved SNR, it does not similarly modulate the subjective experience of that benefit. A possible explanation is that all three speaker groups were broadly aware of the easier listening condition at 4 dB SNR and adjusted their perceived effort ratings accordingly, regardless of whether they were native or non-native speakers. This would reflect a common subjective anchoring to the perceptual saliency of the noise level, rather than a sensitivity to the differential cognitive demands experienced across groups. Supporting evidence for this interpretation comes from Strand et al. (2018), who found that subjective effort ratings were sensitive to gross changes

in acoustic difficulty but less sensitive to individual differences in cognitive processing efficiency, suggesting that self-report measures primarily reflect the perceived salience of the listening challenge rather than the underlying cognitive cost to any particular speaker.

The three-way interaction between RT, SNR, and Group on NASA-TLX scores was not statistically significant [$\chi^2(4) = 4.95, p = .292, \eta_p^2 = 2.30 \times 10^{-3}$] (Table 4.12), indicating that the combined influence of reverberation level, SNR, and linguistic group membership on subjective perceived effort did not differ substantially across speaker groups. Inspection of the specific three-way LMER coefficients confirmed the negligible magnitude of all these effects: Average RT $\times$ 4 dB SNR $\times$ Malayalam ($\beta = 0.18, p = .31$), Maximum RT $\times$ 4 dB SNR $\times$ Malayalam ($\beta = -0.15, p = .30$), Average RT $\times$ 4 dB SNR $\times$ Kannada ($\beta = -0.05, p = .79$), and Maximum RT $\times$ 4 dB SNR $\times$ Kannada ($\beta = -0.01, p = .92$) were all statistically non-significant. This pattern indicates that regardless of whether speakers were native or non-native, and regardless of the specific combination of reverberation and SNR they encountered, their subjective sense of combined acoustic-linguistic workload did not differ across groups in a way that varied jointly with both acoustic variables simultaneously. This finding is particularly striking because the identical three-factor combination produced a statistically significant and moderate three-way interaction on the objective secondary task measure ($\chi^2(2) = 19.38, p < .001, \eta_p^2 = 0.07$) (Table 4.7), demonstrating that the compound demands of Maximum RT, varying SNR, and non-native language background were clearly differentiated at the level of cognitive resource allocation, yet remained perceptually undifferentiated at the level of subjective self-report.

This objective-subjective dissociation in the three-way interaction is theoretically important and can be understood through several complementary mechanisms. First, the NASA-TLX is a retrospective workload rating tool that captures the speaker's appraisal of perceived effort across six dimensions (Hart & Staveland, 1988). By its nature, it integrates

the speaker's affective experience, motivational state, and perceived performance into a summary rating that may not be sensitive to the fine-grained differences in cognitive resource allocation across conditions that the secondary task is designed to capture. When the three factors of RT, SNR, and linguistic group combine, the resulting cognitive demand may be differentially distributed across groups in ways that are too subtle to be consciously reportable, even though they were measurable at the level of secondary task performance. This interpretation is consistent with the FUEL framework (Pichora-Fuller et al., 2016), which distinguishes between the cognitive processing costs of listening and the affective and motivational dimensions of perceived effort, proposing that these two dimensions are influenced by related but partially independent processes. The objective dual-task secondary measure primarily taps the former, while NASA-TLX primarily captures the latter.

Second, a ceiling effect in NASA-TLX ratings among the Hindi group offers a complementary and important explanation for the non-significant three-way interaction. The pairwise group comparison data clearly showed that Hindi speakers reported consistently the highest NASA-TLX scores across all RT and SNR combinations, with comparatively smaller variability in their ratings (Table 4.13). As a consequence, Hindi speakers may not have been able to sense further changes in their workload when both RT and SNR were varied, because they had already (or nearly) reached their maximum perceived effort under easier conditions (such as minimum RT, 4 SNR). The pairwise RT comparisons for the Hindi group at 0 dB SNR were all non-significant (Minimum vs Average: $t = -0.55$, $p = 1.00$, $d = -0.08$; Minimum vs Maximum: $t = -1.22$, $p = .66$, $d = -0.15$), confirming that Hindi speakers showed negligible RT-related variation in their NASA-TLX ratings under 0 dB SNR. This ceiling in mental load reached at better acoustic conditions compresses the score range for the most burdened group, necessarily attenuating any three-way interaction that depends on differential group sensitivity to the combined $RT \times SNR$ manipulation. Lau et

al. (2019) demonstrated this type of subjective anchoring effect empirically. Using subjective listening effort ratings and physiological pupillometry measures (pupil dilation response) during an auditory task, Lau et al. (2019) found that physiological measures of effort showed greater sensitivity to condition changes than self-report ratings. This finding was specifically found in participants whose subjective ratings were already elevated at baseline, supporting the interpretation that ceiling effects constrain the responsiveness of self-report measures among high-effort Hindi speakers in the present study.

Third, the nature of reverberation as a perceptual challenge may itself contribute to the non-significant three-way interaction on NASA-TLX (Table 4.12). Unlike background noise, which is perceptually salient and immediately attributable to an identifiable external source, reverberation produces a gradual, diffuse blurring of the acoustic signal that speakers may not consciously identify as a discrete increase in difficulty. When reverberation combines with noise at the levels examined in the present study, speakers may perceive the resulting difficulty as simply “harder speech” without specifically attributing additional workload to the reverberant component or to their linguistic group membership. Picou et al. (2016) similarly demonstrated that reverberation significantly degraded performance on speech recognition tasks without eliciting corresponding increases in secondary task response times, a finding attributed to the perceptually non-salient nature of reverberant degradation. Although the present study found that reverberation did affect secondary task performance substantially, the parallel finding of a non-significant RT contribution to NASA-TLX three-way interactions is consistent with this broader pattern of limited perceptual conspicuity of reverberant signal degradation in self-report measures.

Fourth, while no directly comparable three-way interaction between RT, SNR, and non-native speaker group on NASA-TLX has been reported in the published literature to date (reflecting the novelty of this specific design), the past research shows that subjective

effort and actual performance often diverge in complex listening conditions. For example, Seitz et al. (2024) found that while subjective effort reports match performance in simple noise conditions but becomes less reliable when multiple challenges are combined. Similarly, Francis et al. (2021) reported that language background language background affects behavioural and physiological effort, but these differences are not always reflected in self-reports, especially in difficult conditions due to ceiling effects (discussed above).

Taken together, the non-significant NASA-TLX interaction does not mean there are no real differences in listening effort. Instead, it suggests that self-report measures may not be sensitive enough to capture subtle differences in complex listening situations. In contrast, the significant three-way interaction in the secondary task shows that these differences do exist at the cognitive level. Supporting this, Giuliani et al. (2021) and Keur-Huizinga et al. (2024) explicitly cautioned against relying on any single measure of listening effort. The present study reinforces this point (interaction was evident in the objective measure, and not subjective one), highlighting that conclusions based on a single measure can be incomplete.

5.6 Relationship Between Second Language Proficiency and Listening Effort Measures

The Spearman correlation analyses revealed strong positive relationships between LEAP-Q scores (first/second language proficiency in Kannada) and both primary and secondary task performance (Figure 4.11), and moderate negative relationships between LEAP-Q scores and NASA-TLX ratings (Figure 4.12).

The strong positive correlations between LEAP-Q scores and primary task performance (r_s ranging from 0.52 to 0.91 across conditions) (Figure 4.11) confirm that higher proficiency in Kannada was systematically associated with better speech recognition accuracy in noise and reverberation. The relationships were particularly strong under adverse acoustic conditions i.e., at 0 dB SNR across all RTs ($r_s \approx 0.90\text{--}0.91$) (Figure 4.11),

suggesting that proficiency becomes an increasingly critical resource as the acoustic environment deteriorates. This is consistent with the findings advocated by Mattys et al. (2012) that higher-level linguistic knowledge serves as a compensatory resource for speech perception under adverse conditions. When bottom-up acoustic cues are degraded, more proficient speakers can rely more effectively on stored lexical and phonological knowledge to infer missing information, maintaining higher accuracy.

The strong positive correlations between LEAP-Q and secondary task performance ($r_s = 0.83\text{--}0.90$) (Figure 4.11) further confirm that higher Kannada proficiency was associated with lower cognitive listening effort, even after accounting for the influence of acoustic conditions. More proficient speakers were able to process the primary speech task more efficiently, leaving greater cognitive capacity available for the secondary task. Borghini & Hazan (2020) found that listening effort, as measured by a dual-task paradigm, decreased as non-native speakers gained more experience and proficiency in English, supporting a direct link between language proficiency and listening effort efficiency. Similarly, Song & Iverson (2018) demonstrated that greater L2 experience reduced neural effort during speech processing, consistent with the cognitive efficiency gains captured in the present data.

The moderate negative correlations between LEAP-Q and NASA-TLX ratings (r_s ranging from -0.35 to -0.48) (Figure 4.12) confirm that higher proficiency was associated with lower perceived workload, though the relationships were weaker than those observed for objective task performance. The moderate rather than strong nature of the correlation between proficiency and subjective effort is meaningful: it reflects the fact that subjective effort ratings are influenced by multiple factors beyond proficiency alone, including individual differences in metacognitive awareness, personality, fatigue, and tolerance for cognitive challenge. Additionally, the LEAP-Q measure, while validated, captures self-

reported proficiency rather than direct behavioural or neurophysiological indices, which may introduce some measurement noise.

The findings on language proficiency and effort are particularly relevant in the Indian linguistic context, where many individuals communicate daily in a second or third language across social, educational, and professional settings (Kulkarni-Joshi, 2019). The present data suggest that even among individuals who are sufficiently proficient in a second language to participate in an experimental speech perception study, variations in proficiency continue to meaningfully predict cognitive listening load. This has implications for educational settings in India, where students often receive instruction in a language that is not their L1. Students with lower proficiency in the language of instruction would be expected to experience greater listening effort, potentially affecting attention, memory encoding, and academic performance.

The strongest correlations observed under adverse acoustic conditions (0 dB SNR, Maximum RT) are particularly noteworthy in this regard. They indicate that the listening effort penalty associated with lower proficiency is amplified precisely in the conditions that characterise real-world Indian classrooms and public spaces — noisy, reverberant environments where acoustic conditions are frequently poor (Das et al., 2019; Jamir et al., 2014). This convergence between the experimental findings and real-world acoustic contexts gives the present results considerable practical relevance and underlines the importance of dual acoustic and linguistic support strategies for non-native speakers in these settings.

5.7 Dissociation Between Objective and Subjective Measures of Listening Effort

One of the cross-cutting themes that emerges from the results is the consistent dissociation between objective listening effort (as measured by the dual-task paradigm

secondary task performance) and subjective listening effort (as measured by NASA-TLX). Reverberation exerted its strongest and most consistent effects on secondary task performance (Table 4.7) with effect sizes that were among the largest in the dataset, yet failed to produce a significant main effect on NASA-TLX scores (Table 4.12). Conversely, speaker group effects were similarly large across both measures, but the SNR effect, while significant across both, produced larger objective than subjective changes.

This dissociation reflects the multidimensional nature of listening effort as conceptualised by Pichora-Fuller et al. (2016) and Francis & Love (2020). The FUEL framework explicitly distinguishes between cognitive processing demands (what the dual-task paradigm captures) and affective-motivational dimensions of effort, including subjective experience of exertion and workload (what the NASA-TLX captures). These two dimensions are influenced by overlapping but not identical sets of factors. The dual-task paradigm captures online resource depletion during listening — a process that may occur relatively automatically and outside conscious awareness, particularly for well-practised speakers. The NASA-TLX, in contrast, captures the speaker's retrospective and conscious judgement of how hard they worked, which is influenced by affective, motivational, and metacognitive factors in addition to actual cognitive load.

Lau et al. (2019) demonstrated precisely this kind of dissociation in a study where physiological measures (pupil dilation) and subjective ratings yielded different patterns across task conditions, with physiological indices reflecting real-time cognitive load changes more sensitively than self-report. Giuliani et al. (2021) reviewed multiple methods of measuring listening effort and noted that the various measures assess different facets of effort and do not necessarily converge. The present findings add to this empirical base by demonstrating, in a multilingual group with varying language background, that

reverberation specifically represents a form of acoustic degradation that elevates objective cognitive load without producing proportionate increases in subjective effort awareness.

This finding carries methodological implications for researchers investigating listening effort in clinical and applied settings. Relying solely on self-report measures such as the NASA-TLX may underestimate the true cognitive burden imposed by reverberant environments, since participants may not subjectively register the extent to which reverberation is depleting their cognitive resources. Similarly, relying solely on dual-task performance may fail to capture the subjective distress and discomfort associated with linguistically challenging listening situations. Keur-Huizinga et al. (2024) and Tednes & Seeman (2017) both advocated for the combined use of objective and subjective measures in listening effort research, a recommendation that is empirically supported by the dissociations observed in the present data.

CHAPTER 6

SUMMARY AND CONCLUSIONS

The present study provides compelling evidence that listening effort in adverse acoustic conditions is shaped by a complex interplay of acoustic factors (signal to noise ratio- SNR, and reverberation time-RT) and linguistic factors (native language status and second language proficiency), and that these factors interact in meaningful ways. Native Kannada speakers consistently demonstrated superior primary task performance and lower cognitive listening effort compared to Malayalam-Kannada non-native speaker, who themselves outperformed Hindi speakers. These findings suggest that speech perception in noise and the associated cognitive load is influenced not only by language nativity but also by the typological relationship between a speaker's first language (L1) and the target language. Increasing RT and decreasing SNR both degraded speech recognition and elevated cognitive listening effort, with the most adverse compound conditions being especially detrimental for non-native speakers. Subjective listening effort was most strongly predicted by speaker group, while the objective dual-task measures were more sensitive to the effects of reverberation, highlighting an important dissociation between these two dimensions of the listening effort construct.

Second language proficiency emerged as a robust correlate of both objective and subjective listening effort, with higher proficiency associated with better recognition performance, lower cognitive load, and reduced perceived workload, particularly under the most adverse acoustic conditions. Together, these findings underscore the need for multidimensional assessment of listening effort in research and clinical practice, and they have direct implications for the design of acoustic environments in India's linguistically

diverse educational and professional settings. They also contribute to the growing theoretical and empirical literature on listening effort by demonstrating that linguistic background and specifically the typological relationship between L1 and the language of communication is a significant and systematically meaningful predictor of cognitive listening demand in multilingual populations.

6.1 Theoretical Implications and Contextual Significance

The current findings contribute to the theoretical understanding of listening effort in several important ways. First, they demonstrate that the construct of listening effort is sensitive not only to acoustic degradation parameters such as SNR and RT, which have received substantial attention in the literature, but also to the linguistic background and language proficiency of the speaker, which have been comparatively less studied, especially in non-western multilingual contexts. The native-language advantage documented here is consistent with the broader theoretical framework of the Ease of Language Understanding (ELU) model (Rönnberg et al., 2013), which predicts that mismatches between degraded acoustic input and stored linguistic representations require more effortful, explicit processing. For non-native speakers, phonological representations of L2 words are less robust and automatic, creating more frequent mismatches and consequently higher listening effort.

Second, the findings extend the existing understanding of the interaction between linguistic and acoustic factors. The significant three-way interactions among RT, SNR, and Group across both primary and secondary tasks demonstrate that the acoustic and linguistic demands on listening effort do not simply add up independently. Instead, they interact in complex ways, with acoustic degradation amplifying the disadvantages of non-native language status and with linguistic proficiency partially buffering against acoustic

challenges. This interactive pattern supports a resource-based view of listening effort in which both bottom-up (acoustic) and top-down (linguistic) factors compete for limited cognitive resources (Peelle, 2018).

Third, the study's focus on Kannada, a Dravidian language, and its comparison of groups differing in typological proximity to Kannada provides a valuable and relatively rare data point in the listening effort literature, which has been dominated by English as the target language. The gradient of performance across Hindi, Malayalam, and Kannada speakers, explicable in terms of typological distance, suggests that cross-linguistic structural similarity is a meaningful and theoretically grounded predictor of L2 listening effort. Future research could systematically operationalise and measure typological distance to test this relationship more directly.

From an applied perspective, the findings are directly relevant to the acoustic design of educational and professional environments in India. Given the documented evidence that Indian classrooms and public spaces frequently suffer from high noise levels and extended reverberation times (Saravanan et al., 2019; Sundaravadhanan et al., 2017), the present data suggest that non-native speakers of the instructional or working language in these environments are at a double disadvantage: they must manage both the inherent cognitive challenges of L2 processing and the additional cognitive load imposed by poor acoustic conditions. The results argue strongly for acoustic improvements such as sound-absorbing materials, optimised room geometry, and supplementary sound amplification systems, particularly in diverse multilingual educational institutions.

6.2 Limitations and Directions for Future Research

While the present study offers several novel contributions, a number of limitations merit consideration. First, the study examined only two SNR conditions (0 and 4 dB) and

three pre-selected RT conditions. The relatively narrow SNR range may have constrained the range of effects observable, particularly for the subjective NASA-TLX measure. Future research should consider a broader range of SNRs, including positive SNRs where intelligibility is largely preserved but listening effort may still differ across groups, to more fully characterise the SNR-effort function.

Second, while the LEAP-Q provided a useful measure of language proficiency, it is a self-report instrument. Direct proficiency measures using standardised language tests would provide a more objective and fine-grained characterisation of the relationship between L2 competence and listening effort. Additionally, the study did not include measures of working memory capacity, which has been theorised and empirically examined as a moderator of the noise-effort relationship (Strand et al., 2018). Future studies might incorporate measures such as the Reading Span Test or the Digit Span to explore whether working memory interacts with language background in predicting listening effort.

Third, the use of a single secondary task (verbal task) limits the scope of the effort inferences that can be drawn. Other secondary tasks, including reaction time or memory tasks might reveal different aspects of cognitive load. Similarly, incorporating physiological measures such as pupillometry alongside the behavioural and subjective measures would enable a more complete picture of the multi-level effort construct.

Fourth, all participants were young adults with normal hearing. Extending this research to older adults or individuals with hearing loss, both groups known to be particularly susceptible to adverse listening conditions (McGarrigle et al., 2014; Nishida et al., 2024) would be valuable, especially given the additional challenge that these individuals face in linguistically diverse environments.

Finally, the present study examined listening effort only under background noise conditions using six-talker babble. Future research could extend the experimental design to

include more complex masker conditions, such as speech babble in both same-language and different-language contexts, to further investigate the influence of informational masking in addition to energetic masking on listening effort among native and non-native speakers. Such investigations would provide deeper insight into how linguistic similarity between target and masker speech modulates cognitive load during speech perception.

REFERENCES

- Abraham, A. K., & Ravishankar, M. S. (2021). A case study of acoustic intervention in classrooms. *Building Acoustics*, 28(4), 293–308.
<https://doi.org/10.1177/1351010X20975765>
- Adank, P., Evans, B. G., Stuart-Smith, J., & Scott, S. K. (2009). Comprehension of Familiar and Unfamiliar Native Accents Under Adverse Listening Conditions. *Journal of Experimental Psychology: Human Perception and Performance*, 35(2), 520–529. <https://doi.org/10.1037/A0013552>
- Alhanbali, S., Dawes, P., Lloyd, S., & Munro, K. J. (2017). Self-Reported Listening-Related Effort and Fatigue in Hearing-Impaired Adults. *Ear and Hearing*, 38(1), e39–e48. <https://doi.org/10.1097/AUD.0000000000000361>
- Alhanbali, S., Dawes, P., Millman, R. E., & Munro, K. J. (2019). Measures of Listening Effort Are Multidimensional. *Ear & Hearing*, 40(5), 1084–1097.
<https://doi.org/10.1097/AUD.0000000000000697>
- Allen, J. B., & Berkley, D. A. (1979). Image method for efficiently simulating small-room acoustics. *The Journal of the Acoustical Society of America*, Vol. 65(Issue 4).
<https://doi.org/10.1121/1.382599>
- Alqattan, D., & Turner, P. (2021). The effect of background noise on speech perception in monolingual and bilingual adults with normal hearing. *Noise and Health*, 23(110), 67–74. https://doi.org/10.4103/nah.nah_55_20
- Aparna, U., & Hemanth, N. (2019). *Effect of native and non- native multi-talker babble on listening effort in Kannada -Malayalam speaking bilinguals* [Master's Dissertation, All India Institute of Speech and Hearing].
<http://192.168.100.26:8080/xmlui/handle/123456789/1661>

- Bates, D., Mächler, M., Bolker, B. M., & Walker, S. C. (2015). Fitting linear mixed-effects models using lme4. *Journal of Statistical Software*, 67(1).
<https://doi.org/10.18637/JSS.V067.I01>
- Borghini, G., & Hazan, V. (2018). Listening effort during sentence processing is increased for non-native listeners: A pupillometry study. *Frontiers in Neuroscience*, 12(MAR).
<https://doi.org/10.3389/fnins.2018.00152>
- Borghini, G., & Hazan, V. (2020). Effects of acoustic and semantic cues on listening effort during native and non-native speech perception. *The Journal of the Acoustical Society of America*, 147(6), 3783–3794. <https://doi.org/10.1121/10.0001126>
- Brännström, K. J., Rudner, M., Carlie, J., Sahlén, B., Gulz, A., Andersson, K., & Johansson, R. (2021). Listening effort and fatigue in native and non-native primary school children. *Journal of Experimental Child Psychology*, 210.
<https://doi.org/10.1016/j.jecp.2021.105203>
- Brooks, M. E., Kristensen, K., van Benthem, K. J., Magnusson, A., Berg, C. W., Nielsen, A., Skaug, H. J., Mächler, M., & Bolker, B. M. (2017). glmmTMB balances speed and flexibility among packages for zero-inflated generalized linear mixed modeling. *R Journal*, 9(2), 378–400. <https://doi.org/10.32614/RJ-2017-066/>
- Brown, V. A., & Strand, J. F. (2019). Noise increases listening effort in normal-hearing young adults, regardless of working memory capacity. *Language, Cognition and Neuroscience*, 34(5). <https://doi.org/10.1080/23273798.2018.1562084>
- Carhart, R., & Jerger, J. F. (1959). Preferred Method For Clinical Determination Of Pure-Tone Thresholds. *Journal of Speech and Hearing Disorders*, 24(4), 330–345.
<https://doi.org/10.1044/jshd.2404.330>
- Chopra, S., Kaur, H., Pandey, R. M., & Nehra, A. (2018). Development of neuropsychological evaluation screening tool: An education-free cognitive screening

- instrument. *Neurology India*, 66(2), 391–399. <https://doi.org/10.4103/0028-3886.227304>,
- Cueille, R., Lavandier, M., & Grimault, N. (2022). Effects of reverberation on speech intelligibility in noise for hearing-impaired listeners. *Royal Society Open Science*, 9(8). <https://doi.org/10.1098/RSOS.210342>
- Das, P., Talukdar, S., Ziaul, S., Das, S., & Pal, S. (2019). Noise mapping and assessing vulnerability in meso level urban environment of Eastern India. *Sustainable Cities and Society*, 46, 101416. <https://doi.org/10.1016/J.SCS.2019.01.001>
- Degeest, S., Keppler, H., & Corthals, P. (2015). The effect of age on listening effort. *Journal of Speech, Language, and Hearing Research*, 58(5), 1592–1600. https://doi.org/10.1044/2015_JSLHR-H-14-0288
- Desjardins, J. L., & Doherty, K. A. (2014). The effect of hearing aid noise reduction on listening effort in hearing-impaired adults. *Ear and Hearing*, 35(6), 600–610. <https://doi.org/10.1097/AUD.0000000000000028>
- Downey, R. I., Petersen, B. J., Mohanathas, N., Campos, J. L., Montero-Odasso, M., Bherer, L., Pichora-Fuller, M. K., Bray, N. W., Burhan, A. M., Camicioli, R., Fraser, S., Liu-Ambrose, T., Lussier, M., Middleton, L. E., Pieruccini-Faria, F., Phillips, N. A., & Li, K. Z. H. (2026). The effect of hearing ability on dual-task performance following multi-domain training in older adults with mild cognitive impairment: findings from the SYNERGIC trial. *Frontiers in Aging Neuroscience*, 17, 1716733. <https://doi.org/10.3389/FNAGI.2025.1716733/FULL>
- Eranna, P. K., Winston, J. S., Rachna, D., & Mahapatra, P. K. (2025). Influence of cross-linguistic exposure on speech perception under informational masking conditions: a cross-sectional study. *The Egyptian Journal of Otolaryngology* 2025 41:1, 41(1), 155-. <https://doi.org/10.1186/S43163-025-00904-5>

- Faul, F., Erdfelder, E., Lang, A. G., & Buchner, A. (2007). G*Power 3: A flexible statistical power analysis program for the social, behavioral, and biomedical sciences. *Behavior Research Methods* 2007 39:2, 39(2), 175–191.
<https://doi.org/10.3758/BF03193146>
- Finitzo-Hieber, T., & Tillman, T. W. (1978). Room acoustics effects on monosyllabic word discrimination ability for normal and hearing-impaired children. *Journal of Speech and Hearing Research*, 21(3), 440–458.
<https://doi.org/10.1044/JSHR.2103.440>;PAGE:STRING:ARTICLE/CHAPTER
- Francis, A. L., Bent, T., Schumaker, J., Love, J., & Silbert, N. (2021). Listener characteristics differentially affect self-reported and physiological measures of effort associated with two challenging listening conditions. *Attention, Perception, & Psychophysics* 2021 83:4, 83(4), 1818–1841. <https://doi.org/10.3758/S13414-020-02195-9>
- Francis, A. L., & Love, J. (2020). Listening effort: Are we measuring cognition or affect, or both? *Wiley Interdisciplinary Reviews: Cognitive Science*, 11(1), e1514.
<https://doi.org/10.1002/WCS.1514>;WEBSITE:WEBSITE:WIRES;WGROU:STRIN
G:PUBLICATION
- Gafoor, S. A., & Uppunda, A. K. (2023). Speech Perception in Noise and Medial Olivocochlear Reflex: Effects of Age, Speech Stimulus, and Response-Related Variables. *JARO: Journal of the Association for Research in Otolaryngology*, 24(6), 619. <https://doi.org/10.1007/S10162-023-00919-W>
- Gagné, J. P., Besser, J., & Lemke, U. (2017). Behavioral assessment of listening effort using a dual-task paradigm: A review. *Trends in Hearing*, 21, 1–25.
<https://doi.org/10.1177/2331216516687287>

- Gagne, J. P., Poirier, M., Audet, A., & Huot-Montmeny, K. (2023). Speech comprehension and listening effort in noise: a comparison of younger and older adults. *Otorhinolaryngology(Italy)*, 73(4), 161–171. <https://doi.org/10.23736/S2724-6302.23.02484-2>
- Gatehouse, S., & Noble, I. (2004). The Speech, Spatial and Qualities of Hearing Scale (SSQ). *International Journal of Audiology*, 43(2), 85–99. <https://doi.org/10.1080/14992020400050014>
- Geetha, C., Shivaraju, K., Kumar, S., Manjula, P., & Pavan, M. (2014). Development and standardisation of the sentence identification test in the Kannada language. *Journal of Hearing Science*, (2014 Vol. 4 · No. 1). <https://doi.org/10.17430/890267>
- Geethakumary, V. (2002). *A Contrastive Analysis of Hindi and Malayalam*. <https://www.languageinindia.com/sep2002/chap5.html#biblio>
- Giuliani, N. P., Brown, C. J., & Wu, Y. H. (2021). Comparisons of the Sensitivity and Reliability of Multiple Measures of Listening Effort. *Ear and Hearing*, 42(2), 465–474. <https://doi.org/10.1097/AUD.0000000000000950>
- Goslin, J., Duffy, H., & Floccia, C. (2012). An ERP investigation of regional and foreign accent processing. *Brain and Language*, 122(2), 92–102. <https://doi.org/10.1016/j.bandl.2012.04.017>
- Gosselin, P. A., & Gagné, J. P. (2011). Older adults expend more listening effort than young adults recognizing speech in noise. *Journal of Speech, Language, and Hearing Research*, 54(3), 944–958. [https://doi.org/10.1044/1092-4388\(2010/10-0069\)](https://doi.org/10.1044/1092-4388(2010/10-0069))
- Hart, S. G., & Staveland, L. E. (1988). Development of NASA-TLX (Task Load Index): Results of Empirical and Theoretical Research. *Advances in Psychology*, 52(C), 139–183. [https://doi.org/10.1016/S0166-4115\(08\)62386-9](https://doi.org/10.1016/S0166-4115(08)62386-9)

- Hartig, F. (2024). *DHARMa: residual diagnostics for hierarchical (multi-level/mixed) regression models*. <https://cran.r-project.org/web/packages/DHARMa/vignettes/DHARMa.html>
- He, S. C., & Lee, W. (2022). Generalized linear mixed-effects models for studies using different sets of stimuli across conditions. *Frontiers in Psychology, 13*, 955722. <https://doi.org/10.3389/FPSYG.2022.955722>
- Hicks, C. B., & Tharpe, A. M. (2002). Listening effort and fatigue in school-age children with and without hearing loss. *Journal of Speech, Language, and Hearing Research : JSLHR, 45*(3), 573–584. [https://doi.org/10.1044/1092-4388\(2002/046\)](https://doi.org/10.1044/1092-4388(2002/046))
- Hiremath, D., & Hemanth, N. (2019). *Effect of enhanced signal on listening effort in older adults with and without hearing loss* [Master's Dissertation, All India Institute of Speech and Hearing]. <http://192.168.100.26:8080/xmlui/handle/123456789/1705>
- Holube, I., Haeder, K., Imbery, C., & Weber, R. (2016). Subjective Listening Effort and Electrodermal Activity in Listening Situations with Reverberation and Noise. *Trends in Hearing, 20*. <https://doi.org/10.1177/2331216516667734>
- Hughes, S. E., Rapport, F., Watkins, A., Boisvert, I., McMahon, C. M., & Hutchings, H. A. (2019). Study protocol for the validation of a new patient-reported outcome measure (PROM) of listening effort in cochlear implantation: The Listening Effort Questionnaire-Cochlear Implant (LEQ-CI). *BMJ Open, 9*(7). <https://doi.org/10.1136/BMJOPEN-2018-028881>,
- Hunter, C. R. (2022). Listening Over Time: Single-Trial Tonic and Phasic Oscillatory Alpha-and Theta-Band Indicators of Listening-Related Fatigue. *Frontiers in Neuroscience, 16*, 915349. <https://doi.org/10.3389/FNINS.2022.915349/TEXT>

- Jain, C., Konadath, S., Vimal, B. M., & Suresh, V. (2014). Influence of Native and Non-Native Multitalker Babble on Speech Recognition in Noise. *Audiology Research*, 4(1), 89. <https://doi.org/10.4081/AUDIORES.2014.89>
- Jamir, L., Nongkynrih, B., & Gupta, S. K. (2014). Community Noise Pollution in Urban India: Need for Public Health Action. *Indian Journal of Community Medicine : Official Publication of Indian Association of Preventive & Social Medicine*, 39(1), 8. <https://doi.org/10.4103/0970-0218.126342>
- Janardhan, L., & Devi, N. (2012). *Effect of Reverberation on Acceptable Noise Level in Individuals with Normal Hearing and Hearing Impairment* [Master's Dissertation, All India Institute of Speech and Hearing]. <http://192.168.100.26:8080/xmlui/handle/123456789/1117>
- Jeong, M., Oh, H., Park, J., & Song, J. (2024). Listening effort under different rates of speech. *The Journal of the Acoustical Society of America*, 155(3_Supplement), A80–A80. <https://doi.org/10.1121/10.0026879>
- Johns, M. A., Calloway, R. C., Karunathilake, I. M. D., Decruy, L. P., Anderson, S., Simon, J. Z., & Kuchinsky, S. E. (2024). Attention Mobilization as a Modulator of Listening Effort: Evidence From Pupillometry. *Trends in Hearing*, 28, 23312165241245240. <https://doi.org/10.1177/23312165241245240>
- Kanigalpula, A., Nisha, K. V., & Pitchaimuthu, A. N. (2025). *EFFECT OF REVERBERATION AND AGE ON SPEECH RECOGNITION AND WORKING MEMORY SKILLS OF SCHOOL GOING CHILDREN* [Master's Dissertation, All India Institute of Speech and Hearing]. <http://203.129.241.86:8080/xmlui/handle/123456789/6012>
- Kanigalpula, A., Nisha, K. V., & Pitchaimuthu, A. N. (2026). Can you understand me in class? Effects of age and reverberation on speech recognition in school- age children.

International Journal of Pediatric Otorhinolaryngology, 203, 112772.

<https://doi.org/10.1016/J.IJPORL.2026.112772>

Keur-Huizinga, L., Kramer, S. E., De Geus, E. J. C., & Zekveld, A. A. (2024). A Multimodal Approach to Measuring Listening Effort: A Systematic Review on the Effects of Auditory Task Demand on Physiological Measures and Their Relationship. *Ear and Hearing*, 45(5), 1089–1106.

<https://doi.org/10.1097/AUD.0000000000001508>,

Kilman, L., Zekveld, A. A., Hällgren, M., & Rönnberg, J. (2015). Subjective ratings of masker disturbance during the perception of native and non-native speech. *Frontiers in Psychology*, 6. <https://doi.org/10.3389/FPSYG.2015.01065>

Kilman, L., Zekveld, A., Hällgren, M., & Rönnberg, J. (2014). The influence of non-native language proficiency on speech perception performance. *Frontiers in Psychology*, 5(JUL), 56676. <https://doi.org/10.3389/FPSYG.2014.00651/BIBTEX>

Kim, S.-E., Goldrick, M., & Bradlow, A. R. (2025). Talker-specific perceptual adaptation to second-language speech-in-noise: Tuning-in to the talker while tuning-out the noise. *The Journal of the Acoustical Society of America*, 157(6), 4184–4195.

<https://doi.org/10.1121/10.0036844>

Klink, K. B., Schulte, M., & Meis, M. (2012). Measuring listening effort in the field of audiology-a literature review of methods, part 1 Messungen von Höranstrengung im Bereich der Audiologie-eine literaturgestützte Methodenübersicht, Teil 1. In *KlinkZ Audiol* (Vol. 51, Number 2).

Koerner, T. K., & Zhang, Y. (2017). Application of Linear Mixed-Effects Models in Human Neuroscience Research: A Comparison with Pearson Correlation in Two Auditory Electrophysiology Studies. *Brain Sciences*, 7(3), 26.

<https://doi.org/10.3390/BRAINSKI7030026>

- Krishnamurti, B. (2003). The Dravidian languages. *The Dravidian Languages*, 1–545.
<https://doi.org/10.1017/CBO9780511486876>
- Kulkarni-Joshi, S. (2019). Linguistic history and language diversity in India: Views and counterinterviews. *Journal of Biosciences*, 44(3). <https://doi.org/10.1007/s12038-019-9879-1>
- Lau, M. K., Hicks, C., Kroll, T., & Zupancic, S. (2019). Effect of Auditory Task Type on Physiological and Subjective Measures of Listening Effort in Individuals With Normal Hearing. *Journal of Speech, Language, and Hearing Research : JSLHR*, 62(5), 1549–1560. https://doi.org/10.1044/2018_JSLHR-H-17-0473
- Lecumberri, M. L. G., Cooke, M., & Cutler, A. (2010). Non-native speech perception in adverse conditions: A review. *Speech Communication*, 52(11–12), 864–886.
<https://doi.org/10.1016/J.SPECOM.2010.08.014>
- Lee, S. M., Park, C. J., & Haan, C. H. (2022). Investigation of the Appropriate Reverberation Time in Learning Spaces for Elderly People Using Speech Intelligibility Tests. *Buildings*, 12(11). <https://doi.org/10.3390/BUILDINGS12111943>
- Levene, H. (1960). *Levene, H. (1960) Robust Tests for Equality of Variances. In Olkin, I., Ed., Contributions to Probability and Statistics, Stanford University Press, Palo Alto, 278-292. - References - Scientific Research Publishing.*
<https://www.scirp.org/reference/referencespapers?referenceid=2363177>
- Liljencrants, J., Lindblom, B., & Lindblom, B. (1972). Numerical Simulation of Vowel Quality Systems: The Role of Perceptual Contrast. *Language*, 48(4), 839.
<https://doi.org/10.2307/411991>
- Maitreyee, R., & Goswami, S. P. (2009). *Language Proficiency Questionnaire: An Adaptation of LEAP-Q in Indian Context* [Master's Dissertation, All India Institute of Speech and Hearing]. <http://192.168.100.26:8080/xmlui/handle/123456789/963>

- Maluf, Y. S., Ferrari, S. L. P., & Queiroz, F. F. (2022). *Robust beta regression through the logit transformation*. <https://doi.org/10.1007/s00184-024-00949-1>
- Manikandan, M., & Geetha, C. (2020). *The Influence of Age, Cognition and Proficiency of the Non-Native Background Language on Speech Recognition of Native Language* [Master's Dissertation, All India Institute of Speech and Hearing].
<http://192.168.100.26:8080/xmlui/handle/123456789/1721>
- Massey, F. J. (1951). The Kolmogorov-Smirnov Test for Goodness of Fit. *Journal of the American Statistical Association*, 46(253), 68–78.
<https://doi.org/10.1080/01621459.1951.10500769>
- Mattys, S. L., Davis, M. H., Bradlow, A. R., & Scott, S. K. (2012). Speech recognition in adverse conditions: A review. *Language and Cognitive Processes*, 27(7–8), 953–978.
<https://doi.org/10.1080/01690965.2012.705006>
- Mayadevi. (1974). *The Development and Standardization of A Common Speech Discrimination Test for Indians* [Master's Dissertation, All India Institute of Speech and Hearing]. <http://192.168.100.26:8080/xmlui/handle/123456789/287>
- McGarrigle, R., Munro, K. J., Dawes, P., Stewart, A. J., Moore, D. R., Barry, J. G., & Amitay, S. (2014). Listening effort and fatigue: What exactly are we measuring? A British Society of Audiology Cognition in Hearing Special Interest Group “white paper.” *International Journal of Audiology*, 53(7), 433–445.
<https://doi.org/10.3109/14992027.2014.890296>
- Meador, D., Flege, J. E., & Mackay, I. R. A. (2000). Factors affecting the recognition of words in a second language. *Bilingualism: Language and Cognition*, 3(1), 55–67.
<https://doi.org/10.1017/S1366728900000134>

- Muhammad, L. N. (2023). Guidelines for repeated measures statistical analysis approaches with basic science research considerations. *The Journal of Clinical Investigation*, 133(11), e171058. <https://doi.org/10.1172/JCI171058>
- Munro, M. J., & Derwing, T. M. (1995). Processing Time, Accent, and Comprehensibility in the Perception of Native and Foreign-Accented Speech. *Language and Speech*, 38(3), 289–306. <https://doi.org/10.1177/002383099503800305>
- Nabelek, A. K., & Pickett, J. M. (1974). Monaural and binaural speech perception through hearing aids under noise and reverberation with normal and hearing-impaired listeners. *Journal of Speech and Hearing Research*, 17(4), 724–739. <https://doi.org/10.1044/JSHR.1704.724>
- Namboodiripad, S. (2016). *UC San Diego UC San Diego Electronic Theses and Dissertations Title An Experimental Approach to Variation and Variability in Constituent Order*. <https://escholarship.org/uc/item/2sv6z8bz>
- Naresh, J. S. (2018). *COMPARISON OF PERFORMANCE ON SPIN-K BETWEEN NON-NATIVE KANNADA SPEAKERS HAVING HINDI AS NATIVE LANGUAGE AND NATIVE KANNADA SPEAKERS*.
- Naresh, J. S., & Yathiraj, A. (2018). *Comparison of performance on SPIN-K between non-native Kannada speakers having Hindi as native language and native Kannada speakers* [Master's Dissertation]. All India Institute of Speech and Hearing.
- Ng, E. H. N., Rudner, M., Lunner, T., & Rönnberg, J. (2013). Relationships between self-report and cognitive measures of hearing aid outcome. *Speech, Language and Hearing*, 16(4), 197–207. <https://doi.org/10.1179/205057113X13782848890774>,
- Nishida, K., Obuchi, C., Shiroma, M., Okamoto, H., & Noguchi, Y. (2024). Effect of background noise and memory load on listening effort of young adults with and

without hearing loss. *Auris Nasus Larynx*, 51(5), 885–891.

<https://doi.org/10.1016/j.anl.2024.08.005>

Ohala, M., & Ohala, J. J. (1992). PHONETIC UNIVERSALS AND HINDI SEGMENT DURATION. *2nd International Conference on Spoken Language Processing, ICSLP 1992*, 831–834. <https://doi.org/10.21437/icslp.1992-272>

Orlikoff, R. F., Schaivetti, N., & Metz, D. E. (2015). *Evaluating Research in Communication Disorders, seventh edition*. 501.

<https://search.worldcat.org/title/862149364>

Pals, C., Sarampalis, A., Hedderik Van Rijn, ;, Başkent, D., Van Rijn, H., & Başkent, D. B. (2015). Validation of a simple response-time measure of listening effort. *The Journal of the Acoustical Society of America*, 138(3), EL187–EL192.

<https://doi.org/10.1121/1.4929614>

Park, H. Y., Viswanathan, N., & Vigeant, M. C. (2025). Evaluating the effects of reverberation on speech intelligibility and head movement using spatial room impulse responses. *The Journal of the Acoustical Society of America*, 157(4_Supplement), A89–A89. <https://doi.org/10.1121/10.0037505>

Peelle, J. E. (2018). Listening Effort: How the Cognitive Consequences of Acoustic Challenge Are Reflected in Brain and Behavior. *Ear and Hearing*, 39(2), 204.

<https://doi.org/10.1097/AUD.0000000000000494>

Peng, Z. E., & Wang, L. M. (2019). Listening effort by native and nonnative listeners due to noise, reverberation, and talker foreign accent during english speech perception.

Journal of Speech, Language, and Hearing Research, 62(4), 1068–1081.

https://doi.org/10.1044/2018_JSLHR-H-17-0423

- Perreau, A. E., Wu, Y. H., Tatge, B., Irwin, D., & Corts, D. (2017). Listening Effort Measured in Adults with Normal Hearing and Cochlear Implants. *Journal of the American Academy of Audiology*, 28(8), 685. <https://doi.org/10.3766/JAAA.16014>
- Petersen, E. B., Wöstmann, M., Obleser, J., Stenfelt, S., & Lunner, T. (2015). Hearing loss impacts neural alpha oscillations under adverse listening conditions. *Frontiers in Psychology*, 6(FEB), 120014. <https://doi.org/10.3389/FPSYG.2015.00177/TEXT>
- Philips, C., Jacquemin, L., Lammers, M. J. W., Mertens, G., Gilles, A., Vanderveken, O. M., & Van Rompaey, V. (2023). Listening effort and fatigue among cochlear implant users: a scoping review. *Frontiers in Neurology*, 14, 1278508. <https://doi.org/10.3389/FNEUR.2023.1278508>
- Pichora-Fuller, M. K., Kramer, S. E., Eckert, M. A., Edwards, B., Hornsby, B. W. Y., Humes, L. E., Lemke, U., Lunner, T., Matthen, M., Mackersie, C. L., Naylor, G., Phillips, N. A., Richter, M., Rudner, M., Sommers, M. S., Tremblay, K. L., & Wingfield, A. (2016). Hearing impairment and cognitive energy: The framework for understanding effortful listening (FUEL). *Ear and Hearing*, 37, 5S-27S. <https://doi.org/10.1097/AUD.0000000000000312>
- Picou, E. M. (2023). Listening Effort Methodologies: Challenges and Future Directions. *Seminars in Hearing*, 44(2), 93. <https://doi.org/10.1055/S-0043-1767668>
- Picou, E. M., Gordon, J., & Ricketts, T. A. (2016). The effects of noise and reverberation on listening effort in adults with normal hearing. *Ear and Hearing*, 37(1). <https://doi.org/10.1097/AUD.0000000000000222>
- Puglisi, G. E., Warzybok, A., Astolfi, A., & Kollmeier, B. (2021). Effect of reverberation and noise type on speech intelligibility in real complex acoustic scenarios. *Building and Environment*, 204, 108137. <https://doi.org/10.1016/J.BUILDENV.2021.108137>

- Rennies, J., Schepker, H., Holube, I., & Kollmeier, B. (2014). Listening effort and speech intelligibility in listening situations affected by noise and reverberation. *The Journal of the Acoustical Society of America*, 136(5), 2642–2653.
<https://doi.org/10.1121/1.4897398>
- Ringbom, H. (2006). *Cross-linguistic Similarity in Foreign Language Learning*.
<https://doi.org/10.21832/9781853599361>
- Romero-Rivas, C., Martin, C. D., & Costa, A. (2015). Processing changes when listening to foreign-accented speech. *Frontiers in Human Neuroscience*, 9(MAR).
<https://doi.org/10.3389/FNHUM.2015.00167>
- Rönnerberg, J., Lunner, T., Zekveld, A., Sörqvist, P., Danielsson, H., Lyxell, B., Dahlström, Ö., Signoret, C., Stenfelt, S., Pichora-Fuller, M. K., & Rudner, M. (2013). The Ease of Language Understanding (ELU) model: Theory, data, and clinical implications. *Frontiers in Systems Neuroscience*, 7(JUNE), 48891.
<https://doi.org/10.3389/FNSYS.2013.00031/BIBTEX>
- Rosenhouse, J., Haik, L., & Kishon-Rabin, L. (2006). Speech perception in adverse listening conditions in Arabic-Hebrew bilinguals. *International Journal of Bilingualism*, 10(2), 119–135.
<https://doi.org/10.1177/13670069060100020101;SUBPAGE:STRING:ABSTRACT;JOURNAL:JOURNAL:IJBA;WEBSITE:WEBSITE:SAGE;REQUESTEDJOURNAL:JOURNAL:IJBA;WGROU:STRING:PUBLICATION>
- Rovetti, J., Sumantry, D., & Russo, F. A. (2023). Exposure to nonnative-accented speech reduces listening effort and improves social judgments of the speaker. *Scientific Reports*, 13(1), 2808. <https://doi.org/10.1038/S41598-023-29082-1>
- Salah, S. M., El-Kabarity, R., El-Kholi, W., & Kamal, N. (2024). Assessment of Listening Effort Experience in Normal and Hearing Impaired Adult Population. *QJM: An*

International Journal of Medicine, 117(Supplement_2).

<https://doi.org/10.1093/QJMED/HCAE175.244>

- Sarampalis, A., Kalluri, S., Edwards, B., & Hafter, E. (2009). Objective measures of listening effort: Effects of background noise and noise reduction. *Journal of Speech, Language, and Hearing Research*, 52(5), 1230–1240. [https://doi.org/10.1044/1092-4388\(2009/08-0111\)](https://doi.org/10.1044/1092-4388(2009/08-0111))
- Saravanan, G., Selvarajan, H. G., & McPherson, B. (2019). Profiling indian classroom listening conditions in schools for children with hearing impairment. *Noise and Health*, 21(99), 83–95. https://doi.org/10.4103/NAH.NAH_15_19
- Schwartz, J. L., Boë, L. J., Vallée, N., & Abry, C. (1997). The Dispersion-Focalization Theory of vowel systems. *Journal of Phonetics*, 25(3), 255–286. <https://doi.org/10.1006/JPHO.1997.0043>
- Seitz, J., Loh, K., & Fels, J. (2024). Listening effort in children and adults in classroom noise. *Scientific Reports*, 14(1), 25200. <https://doi.org/10.1038/S41598-024-76932-7>
- Shastri, U., & Kumar, A. (2015). *Auditory Perceptual Training of Non-native Speakers: Role of Auditory and Cognitive Factors* [Master's Dissertation, All India Institute of Speech and Hearing]. <http://192.168.100.26:8080/xmlui/handle/123456789/62>
- Shetty, H. N., Raju, S., & Singh S, S. (2023). The relationship between age, acceptable noise level, and listening effort in middle-aged and older-aged individuals. *Journal of Otology*, 18(4), 220–229. <https://doi.org/10.1016/j.joto.2023.09.004>
- Simantiraki, O., Wagner, A. E., & Cooke, M. (2023). The impact of speech type on listening effort and intelligibility for native and non-native listeners. *Frontiers in Neuroscience*, 17, 1235911. <https://doi.org/10.3389/FNINS.2023.1235911/FULL>

- Smeds, K., Wolters, F., & Rung, M. (2015). Estimation of Signal-to-Noise Ratios in Realistic Sound Scenarios. *Journal of the American Academy of Audiology*, 26(2), 183–196. <https://doi.org/10.3766/JAAA.26.2.7>
- Song, J., & Iverson, P. (2018). Listening effort during speech perception enhances auditory and lexical processing for non-native listeners and accents. *Cognition*, 179, 163–170. <https://doi.org/10.1016/j.cognition.2018.06.001>
- Spandita, M. (2025). *The Effect of Verbal Working Memory Training on Auditory Processing Abilities, Visuospatial Working Memory and Listening Effort in Children with Central Auditory Processing Disorders* [Doctoral Thesis, All India Institute of Speech and Hearing]. <http://203.129.241.86:8080/xmlui/handle/123456789/7729>
- Strand, J. F., Brown, V. A., Merchant, M. B., Brown, H. E., & Smith, J. (2018). Measuring Listening Effort: Convergent Validity, Sensitivity, and Links With Cognitive and Personality Measures. *Journal of Speech, Language, and Hearing Research*, 61(6), 1463–1486. https://doi.org/10.1044/2018_JSLHR-H-17-0257
- Stringer, L., & Iverson, P. (2020). Non-native speech recognition sentences: A new materials set for non-native speech perception research. *Behavior Research Methods*, 52(2), 561–571. <https://doi.org/10.3758/S13428-019-01251-Z/FIGURES/2>
- Suman M, R., & Geetha, C. (2018). *The Influence of Proficiency of The Non-Native Background Language on Speech Recognition of Native Language* [Master's Dissertation, All India Institute of Speech and Hearing]. <http://192.168.100.26:8080/xmlui/handle/123456789/1518>
- Sundaravadhanan, G., Selvarajan, H., & McPherson, B. (2017). Classroom Listening Conditions in Indian Primary Schools: A Survey of Four Schools. *Noise & Health*, 19(86), 31. <https://doi.org/10.4103/1463-1741.199240>

- Takata, Y., & Nábelek, A. K. (1990). English consonant recognition in noise and in reverberation by Japanese and American listeners. *The Journal of the Acoustical Society of America*, 88(2), 663–666. <https://doi.org/10.1121/1.399769>
- Tednes, M., & Seeman, S. E. (2017). Listening Effort Outcome Measures in Adult Populations. *AuD Capstone Projects - Communication Sciences and Disorders*, 5–11. <https://ir.library.illinoisstate.edu/aucpcsd/4>
- Vaidyanath, R., & Yathiraj, A. (2014). Screening checklist for auditory processing in adults (SCAP-A): Development and preliminary findings. *Journal of Hearing Science*, 4(1), 27–37. <https://www.journalofhearingscience.com/pdf-120593-49255?filename=49255.pdf>
- van Engen, K. J., Chandrasekaran, B., & Smiljanic, R. (2012). Effects of Speech Clarity on Recognition Memory for Spoken Sentences. *PLOS ONE*, 7(9), e43753. <https://doi.org/10.1371/JOURNAL.PONE.0043753>
- van Wijngaarden, S. J., Steeneken, H. J. M., & Houtgast, T. (2002). Quantifying the intelligibility of speech in noise for non-native listeners. *The Journal of the Acoustical Society of America*, 111(4), 1906–1916. <https://doi.org/10.1121/1.1456928>
- Xu, X. S., Samtani, M. N., Dunne, A., Nandy, P., Vermeulen, A., & De Ridder, F. (2013). Mixed-effects beta regression for modeling continuous bounded outcome scores using NONMEM when data are not on the boundaries. *Journal of Pharmacokinetics and Pharmacodynamics*, 40(4), 537–544. <https://doi.org/10.1007/S10928-013-9318-0>
- Yathiraj, A., Jain S.N, & Amruthavarshini B. (2018). *Comparison of performance on SPIN-K between non-native Kannada speakers having Malayalam as native language and native Kannada speakers* [Master's Dissertation, All India Institute of Speech and Hearing]. <http://192.168.100.26:8080/xmlui/handle/123456789/1497>

Yu, Z., Guindani, M., Grieco, S. F., Chen, L., Holmes, T. C., & Xu, X. (2022). Beyond t test and ANOVA: applications of mixed-effects models for more rigorous statistical analysis in neuroscience research. *Neuron*, *110*(1), 21–35.
<https://doi.org/10.1016/J.NEURON.2021.10.030>

APPENDIX A

Institute Review Board (IRB) approval for Bio-Behavioural Research proposal involving human subjects at AIISH

Informed Consent

I have been informed about the aims, objectives and the procedure of the study. The possible risks-benefits of my participation as human participants in the study are clearly understood by me. I understand that I have a right to refuse participation as subject or withdraw my consent at any time without adversely affecting my/my ward's treatment at AIISH. I am also aware that by subjecting to this investigation, I will have to give more time for assessments by the investigating team and that these assessments may not result in any benefits to me. I have the freedom to write to Chairman AEC, in case of any violation of these provisions without the danger of my being denied any rights to secure the clinical services at this institute.

I _____, the undersigned, give my consent to be participant of this investigation/study/program.

Signature of Parent/
Guardian
(Name and Address)

Signature of Witness

(Name of Witness)

Signature of Investigator:

Name and Designation:

Date:

APPENDIX B

Institute Review Board (IRB) approval for Bio-Behavioural Research proposal involving human subjects at AIISH



All India Institute of Speech and Hearing

(An autonomous Institute under the
Ministry of Health and Family Welfare, Govt. of India)
Center of Excellence - Assessed & accredited by NAAC with 'A' Grade
ISO 9001: 2015 Implemented Institute
Manasagangothri, Mysuru - 570 006

ಅಖಿಲ ಭಾರತ ವಾಕ್ ಶ್ರವಣ ಸಂಸ್ಥೆ
ಮಾನಸಗಂಗೋತ್ರಿ, ಮೈಸೂರು - 570 006

अखिल भारतीय वाक् श्रवण संस्थान
मानसगंगोत्री, मैसूर - 570 006

INSTITUTE REVIEW BOARD (IRB) APPROVAL FOR BIO-BEHAVIOURAL RESEARCH PROPOSAL INVOLVING HUMAN SUBJECTS AT AIISH

AIISH Institute Review Board (IRB)

Ref: SH/IRB/M.4/AUD.22/2025-26

Date: 06.01.2026

Title of the Dissertation

: Native versus Non-Native Speech
Perception in Noise: Effect of
Reverberation and Signal to Noise
Ratio on Listening Effort

Name of the Investigator

: Shivani Sharma

Guide

: Dr. Nisha K V

Co-Guide (if any)

: Dr. P. Arivudai Nambi

Proposed Duration of the Project

: 08 months

Source of funding

: Non-funded Master's dissertation

Estimated budget

: Nil

Reference Number of the Project

: Nil

Date on which the IRB meeting was held

: 02.01.2026

A clear statement of the decision reached IRB meeting

(In the vent of the proposal is not approved, a statement
of reasons for the same must be indicated)

: APPROVED

Advice & suggestions (if any)

: NIL

Signature & Name of the Member Secretary-IRB

Dr. Animesh Barman

Professor of Audiology

All India Institute of Speech and Hearing, Mysore-570006

ಸಂಸ್ಥಾಪಕ ಸಮಿತಿ

Institutional Review Board

अखिल भारतीय वाक् श्रवण संस्थान

All India Institute of Speech and Hearing

मानसगंगोत्री / Manasagangothri

ಮೈಸೂರು, Mysuru 570 006.

Phone : 0821-2502000 / 2502100 Toll Free No. 18004255218 Fax : 0821-2510515,
e-mail : director@aiishmysore.in / @aiishmysuruofficial, @aiishmysuru.
website : www.aiishmysore.in

APPENDIX-C

Screening Checklist for Auditory Processing in Adults (SCAP-A) – Modified 2-point rating scale (Report by Self)

Name:

Date:

Age:

Gender: Male/Female

Please tick (✓) the most appropriate option

No.	Questions	Score	
		Present	Absent
1	Do you require frequent repetitions while listening to someone who does not have a speech problem?	1	0
2	Can you pay attention to someone speaking continuously for more than 10 minutes? E.g. Listening to a conversation.	0	1
3	Do you find it difficult to attend to speech in the presence of background noise? E.g. Television at normal volume/fan at high speed.	1	0
4	Do you have trouble recalling what was said in the correct order? E.g. 5 different (non-routine) things in the order you have done them	1	0
5	Do you forget what was told to you within a short span of time (within a minute)? E.g. To buy a particular item from a shop	1	0
6	Do you have difficulty in understanding speech in the presence of background noise (when the television/fan at full speed)?	1	0
7	Can you recall the names of 5 of your school/ college friends, who you have not met after you left school/college?	0	1
8	Have you been told that you take longer than others to respond when your friends or family talk to you?	1	0
9	Do you have difficulty in responding to two people talking at the same time? E.g. In a group, when two people answer/ask a question at the same time.	1	0
10	Do you feel it is difficult to understand someone's speech when you cannot see his or her face? E.g. When the person's face is turned away from you.	1	0
11	Do you have difficulty in remembering numbers, especially telephone/vehicle/door numbers, bus numbers, account numbers?	1	0
12	Do others report that you do not attend to them when they suddenly start talking to you?	1	0

Thank you for the careful completion of this questionnaire

^aThe scoring for the checklist was modified into a 2-point rating based on the results obtained.

APPENDIX-D

Neuropsychological Evaluation Screening Tool (NEST)

Neuropsychological Evaluation Screening Tool (NEST)				 Sakshi Chopra, Harshdeep Kaur & Ashima Nehra Clinical Neuropsychology, Neuroscience Center, All India Institute of Medical Sciences, New Delhi (© 2015, All Rights Reserved)	
Name:		Education:		Diagnosis/ Provisional Diagnosis:	
Age:		Handedness:			
Sex: ☐ M ☐ F ☐		Date:			
Attention, Immediate Recall					
1	"I will tell you a name & an address. I want you to remember it & repeat after me. I will ask you the same again after sometime." "एक नाम और पता बोला जाएगा, ध्यान से सुनिए और बाद में दोहराएं। आपको दो कुछ समय बाद पूछा जाएगा। मेरी बात कहिए।"	Suraj Ji, Phool Mandi, Nai Sadak सुरज जी, फूल मंडी, नई सड़क	If all correct, Score = 0 If any error, Score = 1	Score:	
Executive Function, Verbal Fluency					
2	Name all the things we eat in 1 minute. You can name fruits, vegetables, dals/pulses, grains, etc. "1 मिनट में आप जिसकी भी खाने की चीजों का नाम बता सकते हैं, उलने बतलाएं। आप कसो, सब्जियों, दालों, मीठ, आदि के नाम से सकते हैं।"		If ≥ 18 items, Score = 0 If < 18 items, Score = 1	Score:	
Executive Function - Abstraction					
3	"What is the similarity between wood & paper?" "क्या आप लकड़ी और कागज में एक समानता बता सकते हैं?"	Can be burnt; Come from trees; Can be used to write जल सकते हैं पेड़ों से आते हैं लिखा जा सकता है	Any one similarity correctly identified, Score = 0 Error = 1	Score:	
Visuo-spatial, Executive, Temporal Sequencing, Attention					
4	"Here is one finger (point towards the finger), one dot (point towards the dot), then two fingers (point), & then two dots (point). I want you to join them with a pencil." Elaborate - "One finger is to be joined with one dot, one dot with two fingers & two fingers with two dots; i.e. in an increasing order." (Let the patient join them). Correct the patient if they go wrong. "Are the instructions clear? Now you have to join them from 1 to 5." "यहाँ है एक उंगली (उंगली दिखाएँ), एक बिंदु (बिंदु दिखाएँ), दो उंगलियाँ (दिखाएँ) और दो बिंदु (दिखाएँ)। आपको पेन्सिल के साथ एक उंगली को एक बिंदु से जोड़ना है, एक बिंदु को दो उंगलियों से, और दो उंगलियों को दो बिंदु से (दोनों को जोड़ने दो) उंगलियों और बिंदु को बढ़ते हुए क्रम में जोड़ना है। (दोनों को अलग-अलग करने दो) अगर सही देखल करें, तो उससे ठीक करें, "अब आपको ऐसे ही 1 से 5 तक उंगलियाँ और बिंदु जोड़ने होंगे।"		If all fingers & dots joined in the correct sequence, Score = 0 If any error, Score = 1	Score:	
Example: 					
Please turn to continue here					

Visuoconstructional Ability			
5	"Please copy the design below"	If angles & lines drawn properly & all lines joined, Score = 0 If any error, Score=1	Score:
	"जैसा चित्र बना है, वैसा ही बना देखिए"		
Delayed Recall			
6	"Can you tell me the name & address I had asked you to remember?"	If any error, Score=1 If all correctly named, Score=0	Score:
	"क्या आप मुझे वो नाम और पता बता सकते हैं जो मैंने पहले बताया था?"		
Obtained Score = /6		Interpretation ☐ ≥ 3 Errors, cognitive impairment likely ☐ ≤ 2 Errors, cognitive impairment is unlikely	

APPENDIX -E

MODIFIED LANGUAGE PROFICIENCY QUESTIONNAIRE

Developed by Yathiraj A., Jain S.N., and Amruthavarshini B. (2018)

Further Modified for the Present Study — Reading and Writing Components Excluded

Name:

Age:

Gender: Male / Female

Instructions: Please read the below given information carefully and choose the most appropriate choice. Respond to all eight points by either filling in blanks or ticking (✓) the most appropriate response. (Note: L1 refers to the first language that you learnt; L2 refers to the second language that you learnt; L3 refers to the third language that you learnt)

1. Name all the languages you have learnt since your childhood in the order of acquisition of the language.

Order of Languages Acquired	Language Name
L1	
L2	
L3	

2. Since when have you been using your L1, L2 and L3 for understanding and speaking? (*Note. Reading and writing components are not assessed in this modified version. Please tick (✓) one duration per language for understanding & speaking*)

Duration (in years)	Understanding			Speaking		
	L1	L2	L3	L1	L2	L3
Less than 5 years						
5 to 10 years						
10.1 to 15 years						
Greater than 15 years						

3. How would you mark your level of proficiency for understanding and speaking? (*Note. Reading and writing components are not assessed in this modified version. Please tick (✓) one level per language for understanding & speaking*)

Level of Proficiency	Understanding			Speaking		
	L1	L2	L3	L1	L2	L3

Low proficiency						
Fair proficiency						
Good proficiency						
Native like / Perfect proficiency						

4. How would you rate your ability to switch between the languages? (Note. Please tick (✓) one of the ratings)

Rating Scale	Response (✓)
Low Ability	
Fair Ability	
Good Ability	
Perfect Ability	

5. Please tick (✓) which language you use maximum for the below mentioned situations: (Note. Please tick (✓) one language per situation)

Sl. No.	Situations	L1	L2	L3
a	Interaction with family			
b	Education / work			
c	Listening to instruction tapes at school			
d	Text books			
e	Dictionary			
f	Story books			
g	Newspapers			
h	Internet source			
i	Writing			
j	Interacting with friends			
k	Interacting with neighbours			
l	Watching TV / YouTube			
m	Listening to the radio (music)			
n	Market places			

6. On a scale of one to four, how often do you use the languages known to you in the following situations? (Rating key: 1 = never; 2 = Sometimes; 3 = Most of the time; 4 = Always; Note. Please write the numbers 1, 2, 3, or 4, for each situation per language)

Sl. No.	Situations	L1	L2	L3
---------	------------	----	----	----

a	Interaction with family			
b	Education / work			
c	Listening to instruction tapes at school			
d	Text books			
e	Dictionary			
f	Story books			
g	Newspapers			
h	Internet source			
i	Writing			
j	Interacting with friends			
k	Interacting with neighbours			
l	Watching TV / YouTube			
m	Listening to the radio (music)			
n	Market places			

7. How frequently do others identify you as a native speaker based on your accent or pronunciation in the language? *(Note. Please tick (✓) one rating per language)*

Rating Scale	L1	L2	L3
Never			
Sometimes			
Most of the time			
Always			

8. For how many hours do you use the following languages? *(Note. Please tick (✓) one duration per language)*

Duration	L1	L2	L3
Greater than 2 hours			
Greater than 3 hours			
Greater than 4 hours			
Greater than 5 hours			

Note: Refer Scoring key for analysis

SCORING KEY

MODIFIED LANGUAGE PROFICIENCY QUESTIONNAIRE Developed by Yathiraj A., Jain S.N., and Amruthavarshini B. (2018) (Reading and Writing components excluded)

Instructions to professional scoring the scale: Please score the responses on a scale of 1 to 4 for each skill / question as directed.

1. Name all the languages you have learnt since your childhood in the order of acquisition of the languages.

No score (Information to be used for descriptive analysis)

2. Since when have you been using your L1, L2 and L3 for understanding and speaking?

Duration (in years)	Scores	Understanding			Speaking			
		L1	L2	L3	L1	L2	L3	
Less than 5 yrs	1							
5 to 10 yrs	2							
10.1 to 15 yrs	3							
Greater than 15 yrs	4							
Total Scores		L1 =/8			L2 =/8			L3 = /8

3. How would you mark your level of proficiency for understanding and speaking?

Level of Proficiency	Scores	Understanding			Speaking			
		L1	L2	L3	L1	L2	L3	
Low proficiency	1							
Fair proficiency	2							
Good proficiency	3							
Native like/ Perfect proficiency	4							
Total Scores		L1 =/8			L2 =/8			L3 =/8

4. How would you rate your ability to switch between the languages?

Rating Scale	Scores	Response			
Low Ability	1				
Fair Ability	2				
Good Ability	3				
Perfect Ability	4				

Total Scores	L1 =	/4	L2 =	/4
---------------------	-------------	-----------	-------------	-----------

5. Tick (✓) which language you use maximum for the following situations: **No score** (Information to be used for descriptive analysis)

6. On a scale of one to four, how often do you use the languages known to you in the following situations? (Instruction to professional scoring the scale: Total the ratings given per language)

Sl. No.	Situations	L1	L2	L3
a	Interaction with family			
b	Education / work			
c	Listening to instruction tapes at school			
d	Text books			
e	Dictionary			
f	Story books			
g	Newspapers			
h	Internet source			
i	Writing			
j	Interacting with friends			
k	Interacting with neighbours			
l	Watching TV / YouTube			
m	Listening to the radio (music)			
n	Market places			
Total Score		/56	/56	/56

7. How frequently others identify you as a native speaker based on your accent or pronunciation in the language?

Rating Scale	Scores	L1	L2	L3	
Never	1				
Sometimes	2				
Most of the time	3				
Always	4				
Total Score		/4	/4	/4	

8. For how many hours do you use the following languages?

Duration	Scores	L1	L2	L3	
Greater than 2 hours	1				

Greater than 3 hours	2				
Greater than 4 hours	3				
Greater than 5 hours	4				
Total Score		/4	/4	/4	

L1	/84
L2	/84
L3	/84

Grand Total (Maximum Score = 84)

Inclusion cut-off: Native Kannada speakers — minimum score of **50/84**; Non-native Kannada speakers (Hindi L1/ Malayalam L1) — minimum score of **30/84**

Note: Scores are computed for the Kannada language column only (L2 for non-native speakers; L1 for native Kannada speakers). Reading and writing components were excluded from Q2 and Q3 as they are not relevant to the auditory-verbal listening effort paradigm used in the present study. Maximum total score = 84 (Understanding+Speaking duration = 8; Understanding+Speaking proficiency = 8; Language switching = 4; Frequency of use = 56; Accent identification = 4; Daily exposure = 4).

APPENDIX F

MODIFIED NASA TASK LOAD INDEX (NASA-TLX)

Hart, S. G., and Staveland, L. E. (1988)

Name: _____ Age: _____

Gender: Male / Female

Language (L1): _____

Instructions: This questionnaire measures your experience during the listening task you just completed. For each of the six questions below, please rate your experience by circling or marking the number that best describes how you felt. Each item is rated on a scale from **0** (Very Low) to **5** (Very High). (**Note:** There are no right or wrong answers — please respond honestly based on your experience during the task.)

0	1	2	3	4	5
Very Low	Low	Somewhat Low	Somewhat High	High	Very High

1. Mental Demand

"How mentally demanding was the task?"

0	1	2	3	4	5
Very Low	Low	Somewhat Low	Somewhat High	High	Very High

2. Physical Demand

"How physically demanding was the task?"

0	1	2	3	4	5
Very Low	Low	Somewhat Low	Somewhat High	High	Very High

3. Temporal Demand

"How hurried or rushed was the pace of the task?"

0	1	2	3	4	5
Very Low	Low	Somewhat Low	Somewhat High	High	Very High

4. Effort

"How hard did you have to work to accomplish your level of performance?"

0	1	2	3	4	5
Very Low	Low	Somewhat Low	Somewhat High	High	Very High

5. Frustration

"How insecure, discouraged, irritated, stressed, or annoyed were you?"

0	1	2	3	4	5
Very Low	Low	Somewhat Low	Somewhat High	High	Very High

6. Performance

"How successful were you in accomplishing the task?"

(Note: For Performance, a lower rating indicates better performance.)

0	1	2	3	4	5
Very Low	Low	Somewhat Low	Somewhat High	High	Very High

APPENDIX G

```

=====
MODEL COMPARISON (DHARMA + AIC)
=====

---- BETA MODEL ----
Uniformity p: 9.564234e-60
Dispersion p: 0
Outliers p: 0.905703

---- GAMMA MODEL ----
Uniformity p: 2.430831e-33
Dispersion p: 0.016
Outliers p: 2.523107e-05

---- AIC and BIC ----
              df      AIC
model_beta   20 -7238.308
model_gamma  20 12216.166
              df      BIC
model_beta   20 -7123.935
model_gamma  20 12330.540

-----LMER MODEL-----
  add uniformity, dispersion, etc here too from dharma
AIC(model_lmer1)
[1] 5176.694
BIC(model_lmer1)
[1] 5259.955

```