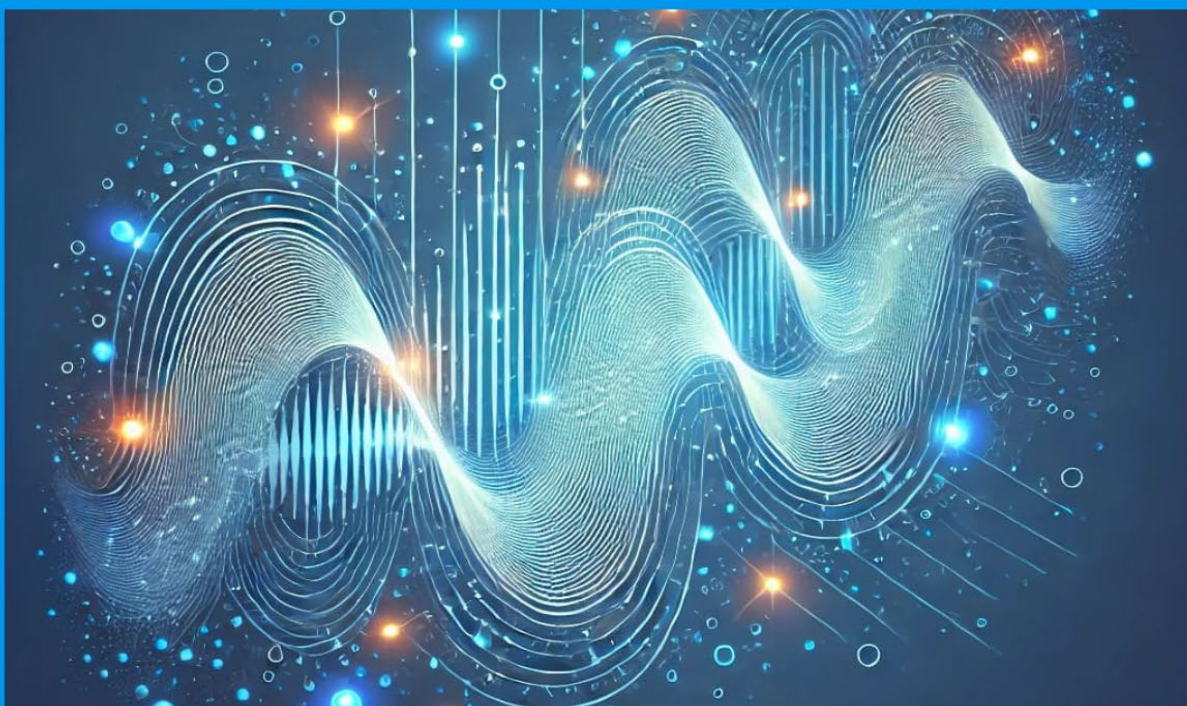


# ਪੰਜਾਬੀ ਵਾਕ ਸਮਰੇਖਿਤ ਵਾਕ ਕੋਰਪਸ



**PUNJABI SENTENCE ALIGNED  
SPEECH CORPUS**

# PUNJABI SENTENCE ALIGNED SPEECH CORPUS



34

*Annotated, quality language data (both-text & speech) and tools in Indian Languages to Individuals, Institutions and Industry for Research & Development - Created in-house, through outsourcing and acquisition.*

**Linguistic Data Consortium for Indian Languages  
Central Institute of Indian Language  
Mysuru, India-570006**

CENTRAL INSTITUTE OF INDIAN LANGUAGES  
Manasagangotri, Mysuru, Karnataka, India, 570006  
www.ciil.org

Punjabi Sentence Aligned Speech Corpus

Dr. Shalinder Singh, Rajesha N., Manasa G., Stephen Fernandes, Nithin S., Roopashri M. R., Dr.  
Narayan Kumar Choudhary, Prof. Shailendra Mohan

*e-ISBN: 978-93-48633-69-9*

*CIIL Publication No.: 1505*

*First published: AD 2025 March :: Phalguna 1946 Śaka*

© Central Institute of Indian Languages, Mysuru 2025

*Publisher:* Prof. Shailendra Mohan, Director, CIIL

*Printing & Publication Unit*

Aleendra Brahma, Officer-in-Charge, Printing Press & Publication Unit

B Manju Bhargavi, Unit-in-Charge, Printing Press & Publication Unit

M.N. Chandrashekar, Compositor, Printing Press & Publication Unit

*Cover design:* Umesh Chamling Rai

## TABLE OF CONTENTS

|                                                      |    |
|------------------------------------------------------|----|
| Table of Contents .....                              | iv |
| Figures .....                                        | iv |
| Tables .....                                         | iv |
| 1 Speech Annotation .....                            | 1  |
| 1.1 Introduction .....                               | 1  |
| 1.2 LDC-IL Speech Annotation .....                   | 1  |
| 1.3 Speech Annotation quality assurance .....        | 4  |
| 2 Punjabi Speech Annotation .....                    | 5  |
| 1.4 Overview of Sentence Aligned Speech Corpus ..... | 5  |
| 1.5 Observations .....                               | 5  |
| 1.6 Summary of the Corpus .....                      | 8  |
| 1.7 Summary of Speakers .....                        | 11 |
| 1.8 References .....                                 | 12 |

## FIGURES

|                                                                                      |    |
|--------------------------------------------------------------------------------------|----|
| Figure 1: Gender-wise Distribution of Punjabi Corpus .....                           | 8  |
| Figure 2: Age-wise Distribution of Punjabi Corpus .....                              | 9  |
| Figure 3: Content Type-wise Distribution of Punjabi Corpus .....                     | 9  |
| Figure 4: Gender Distribution in different Content Types of Punjabi Corpus .....     | 10 |
| Figure 5: Gender Age Distribution in different Content Types of Punjabi Corpus ..... | 10 |

## TABLES

|                                                                                |    |
|--------------------------------------------------------------------------------|----|
| Table 1: Representation of Punjabi Sentence Aligned Speech Data Duration ..... | 11 |
| Table 2: Distribution of Speakers of Sentence Aligned Speech Data .....        | 11 |

# 1 SPEECH ANNOTATION

*Narayan Kumar Choudhary, Rejitha K .S.*

## 1.1 INTRODUCTION

Linguistic Data Consortium for Indian Languages (LDC-IL) has been acting as the repository of Indian language datasets since 2008. It has so far released 42 datasets covering 20 scheduled languages. The released datasets include raw text corpora and raw speech corpora. More details about the scheme and the datasets are available in Choudhary, 2021 and Choudhary and Rao, 2020.

This paper presents a documentation of sentence aligned speech corpora in various languages. This is in continuation of the earlier raw speech corpora in various languages. The difference between the raw speech corpora and the sentence aligned corpora can be easily guessed. As the title of the dataset suggests, the sentence aligned corpus is split at the ‘sentence’ or a meaningful ‘utterance’ boundary while the raw speech corpus had speech segments that were usually larger in size.

## 1.2 LDC-IL SPEECH ANNOTATION

LDC-IL speech corpora are collected using two methods (complete details about the raw speech corpus and the process of its creation are included in the raw speech corpus documentation of respective languages).

- a. Read Speech: A vast majority of the speech dataset in LDC-IL is created by reading texts from various sources. These sources are usually part of the corresponding language raw text corpus.
- b. Spontaneous Conversation: These have been used for data collection in some of the languages and conversational data for other languages will be collected in subsequent phases.

We have manually transcribed both of these kinds of data and aligned the transcriptions at the sentence level. Even though the read speech data is created by reading the texts, a manual transcription was necessitated because a lot of times speakers do not reproduce the texts verbatim and so the original texts do not necessarily correspond to the speech recordings. Moreover, the speech transcriptions represent the different variant forms in speech as well as other speech properties (viz false start, mispronunciation etc.) and non-speech properties (viz background noise) - all of these are meticulously marked during the annotation process.

LDC-IL speech data is annotated with the official script of a language, wherever such a script is specified. For example, Malayalam speech data is annotated in Malayalam Script in UTF-8 encoding; Hindi is annotated in Devanagari script and so forth. The languages which do not have an official script are annotated using the most commonly used script for writing that language. All annotations and transcriptions are carried out using the Praat software. The annotators were

provided with the audio files aligned with the corresponding text file for annotation of the read speech data - this greatly reduced the need of annotating from the scratch for such data.

We have provided the annotations at two levels -

- a. Phonetically Normalised Speech Annotation
- b. Orthographically Normalised Speech Annotation

These two levels are discussed in the following subsections below.

### 1.2.1 Phonetically Normalized Speech Annotation

In phonetically normalized speech annotation layer, we provide the narrow transcriptions of speech. At this layer, even though the script used for transcribing speech is the official script of the language, we use non-standard and even non-existent spellings to represent the speech exactly as it is spoken. The detailed guidelines used by the annotators for transcription and annotation at this layer are given in the following subsection.

#### 1.2.1.1 Guidelines for Phonetically Normalized Speech Annotation

Annotation of this aligned data should be carried out as per the pronunciation of the speaker in the audio. The wrong pronunciation, that is, deviation from the text should be transcribed accordingly.

Following are the instances which should be marked in the annotation.

1. Use of ‘#’

‘#’ symbol is used to mark excessively noisy or unnecessary parts of the audio - generally these portions also contain human speech but are not legible. The audio is post-processed to remove such portions from the audio files that are being released publicly. It is used within the sentence.

2. Use of ‘0’

‘0’ is used to mark silences at the beginning and end of the audio files as well as long silences or non-speech sounds in the intermediate portions of the audio files - these are marked only for those portions which do not have any kind of human speech. These portions are also removed from the publicly released audio files in the post-processing step.

3. Silences are removed.

Any silence longer than 50 ms is marked using #.

4. Cut-off speech and intended speech is marked.

Example: [mini]\*ster —shows that the speaker intended to speak minister but spoke mini in an unclear fashion and ster clearly.

5. Annotation of speech disfluency

Restarts/false starts should be marked. For example, if the speaker intends to speak “bengaluru” but speaks “be bengaluru”, this should be marked as be-bengaluru.

## 6. Numbers

All number sequences should be spelled out. Years should be transcribed in spoken format.

## 7. Mispronunciations

If a speaker mispronounces a word and the mispronunciation is not an actual word, transcription should be done as the word is spoken.

## 8. Utterances longer than 30 seconds should be further split into multiple parts - this split is made at the point of a long silence of around 500 ms.

### 1.2.2 Orthographically Normalized Speech Annotation

In Orthographically Normalized Speech Annotation, a broad transcription is given. This implies that the standardised spelling is used for transcription and speech segments such as cut-off speech, intended speech, repetition etc. are not marked. The complete guidelines are reproduced in the following subsection.

#### 1.2.2.1 Guidelines for Orthographically Normalized Speech Annotation

Orthographically normalized textual layer is prepared over the phonetically normalized text by using the guidelines given below

1. Small portion of a word like a grammatical element or a letter is missed then it is corrected as per the correct writing pattern.  
For example, if the speaker speaks ‘avan viitti poyi’ ‘He went home’ in an informal way then it is annotated as ‘avan viittil poyi’ in the proper standard Malayalam sentence. There is no valid word ‘viitti’ in the language, so it is annotated as a proper word according to the context.
2. Any deviation in the phonetically annotated text is corrected according to standard writing form.  
For example, if the audio is “Maiyam Engineering” it must be corrected as “Marine Engineering”.
3. Restarts/false starts are removed. For example, the speaker intends to speak “keralam” but speaks “ke keralam”. If the word “ke” is a valid word morpheme, then it has been kept, otherwise “ke” is not marked.
4. Cut-off speech is written as standard form and remove [ ]\* symbol
5. Incomplete sentence is standardized to the extent of available audio

**Note:** Indian English - Bengali Variant Speech Annotation and Indian English - Kannada Variant Speech Annotation are following a separate guideline which is available in its respective sections.

### 1.3 SPEECH ANNOTATION QUALITY ASSURANCE

Transcriptions are susceptible to human errors, such as typos or spelling mistakes. To ensure data accuracy, both phonetically normalized, and orthographically normalized textual layers undergo third party evaluation. The third-party evaluator must check the transcription against its corresponding audio to rectify any mistakes. The linguist may accept or reject the corrections by means of matching the original and evaluated transcriptions against the audio. The linguist's specialized knowledge is essential in bringing the transcriptions back in line with accuracy. It is ensured that the mistakes encountered by a third party evaluator are also corrected by arbitrating it further by a third linguist.

## 2 PUNJABI SPEECH ANNOTATION

*Dr. Shalinder Singh, Narayan Kumar Choudhary*

### 1.4 OVERVIEW OF SENTENCE ALIGNED SPEECH CORPUS

Punjabi Sentence Aligned Speech Corpus is created by annotating the speech data collected by LDC-IL. A detailed explanation of the Punjabi Raw Speech Corpus will be available in the [Punjabi Raw Speech Corpus](#) (Ramamurthy, L. et. Al, 2019). LDC-IL Punjabi Sentence Aligned Speech files contain an audio file and two textual layers in Gurumukhi script. Each File is named in accordance with its metadata information like language name, speaker id, content id, gender, age, content type etc.

A Typical LDC-IL naming convention for Sentence Aligned Speech data is shown below.

‘Punjabi\_Female\_16To20\_Contemporary Text-T1\_SP-0031\_T1-0001.wav’

LDC-IL Sentence Aligned Speech corpus for Punjabi contains read speech from four content types viz. contemporary text, creative text, sentences and date format. The contemporary text and creative text are sampled from news and essays/novels respectively. The sentences are a collection of phonetically balanced sentence lists - each speaker has typically recorded 25 sentences randomly selected from this set. Date format is kept as uttered by the speaker. The corpus consists of an audio file for each recording and corresponding textual layer consisting of the phonetically normalised and the orthographically normalised annotation.

### 1.5 OBSERVATIONS

LDC-IL sentence level speech annotation strictly follows what the speaker pronounces rather than what is in the prompt sheet. The text has been written in the respective language script and the speech is transcribed as much as the script supports. Two or more different pronunciations can be uttered by the same or different speaker for the same word.

The reading speed differs from reader to reader. Fast reading informants pose difficulty in annotation. Since news items contain sports news, it includes the informant reading all types of numbers. Speakers sometimes wrongly uttered large digit numbers like thousands or lakhs, decimal numbers, fractions etc. It is observed that speakers read Cricket score, Tennis score etc. in their own way and very few speakers read it properly. Most of the speakers show difficulty in pronouncing foreign names (other than native language names) which is frequently appearing in sports and political news. Abbreviations and rarely used words interrupt the reader’s fluency. All these factors contribute to the complexity in speech which makes it a rather difficult task. Since the dialect of annotator can differ from that of the informant, the annotation process may need repetitive hearing in some cases.

The annotation has to discard the data in particular places where the investigator has communicated with the informant. Some background noise like the sound of a bell, bus horn, other people's conversation, baby crying etc. can be heard in the recording. Since this can be heard along with the voice of the informant, they have to be retained. This slows down the annotation process. Vocal noise of informants like coughing, sneezing etc. can also be observed.

### 1.5.1 PHONETIC ALTERNATION IN PUNJABI SPEECH DATA

Reading speech has in-fluency like unwanted pauses, elongated syllables, word fragments, self-corrections, and repetitive words. When speakers notice what they utter then they suspend their speech and add, delete, or replace words they have already produced. Some fluctuated occurrences were detailed as follows:

#### a. Repetition of words

While reading, if the informant observes that the word has been pronounced incorrectly or not in an effective manner then the speaker normally repeats fragments of the word or sometimes the whole word or the phrase. Sometimes the speaker struggles to read the text and repeats when the content is a bit unfamiliar or there are many foreign words which are difficult to pronounce.

#### b. False start

False start is a common phenomenon in most of the speakers and some speakers it is frequent. Usually, it is the repetition of the first word or syllable of the word but sometimes speakers start with some other letter as well.

E.g.: *b - vid ja;*      *əmbərs - əmrɪntəsər;*      *roj - svərojəgar*

#### c. Intended speech

Intended speech occurred when the speaker slow down or fasten up their speech. Typically it happened at the end of the sentence or stopping the reading. But it can be observed in the middle of the sentence too. E.g.: *da:ktər nei 'mɛnu: [a:'ra:m]\* kə'rən dɪ səla:h dɪt̪t̪i hɛ.* Here *[a:'ra:m]\** is not properly audiable but native speaker could easily understood the word because of language proficiency. In the long words the middle of the syllable or phone might not be audible to the listener or skip by the speaker.

#### d. Addition and Deletion

An extra vowel or a consonant or a syllable is added into a word. The letter which is existing in the word or different letter might be added into the word.

E.g.: *soka: sokka:*      *kɪkət- kɪk:ət.*

Deleting a vowel or a consonant or a syllable from a word is called deletion or elision. It is a common phenomenon when a natural language speaker speaks indistinctly.

E.g.: *məɦɪlə - 'me:lə;*      *vjəvəɦa:rə - be:va:rə*

### e. Assimilation and Dissimilation

Speech is a continuous syllabic fragment, so the articulatory organs influence the preceding or following sound. Consonant or vowel is changed to a similar sound because of the influence of a nearby speech segment called assimilation. E.g. *ʈʰəkkər* - *ʈʰəkkə*

Dissimilation is dropping out a syllable or a letter by the influence of adjacent speech segments. E.g: *təhə:ɖa:* - *θo:ɖa:*

### f. Colloquial usage

Some of the speakers have pronounced colloquial forms instead of the standardized form written in the prompt sheet.

E.g.: *ɖo:kʈər* - *ɖa:kʈər*, *vi:kə:s* - *bi:kə:s*

### g. Lengthening and Shortening

Short and long vowels are interchanged in the recordings at several places.

E.g.: Lengthening: *məha:n* - *mə:ha:n*  
Shortening: *əɖʰi:ja:ɖmək* - *əɖʰi:ja:ɖmæk*

### h. Substandard and Shortening

It has been observed that some speakers have consistently replaced the aspirated sounds with their unaspirated counterparts.

E.g.: *utsa:ha* > *usta:ha*; *viɖʱajo:tsava* > *viɖʱajo:stava*;

### i. Phone variation

It is the alternative pronunciation of the word, and which does not affect the meaning. Both pronunciations are considered to be in free variations.

E.g. *kəfən* - *kəp ən*, *məzil* - *məɳil*, *məcc əɾ* - *məcc əɽ*

### j. Final vowel modification

In continuous speech the final vowel gets modified at times in some of the speakers:

E.g. *ra rəpatl* - *ra rəpət i:*

## 1.6 SUMMARY OF THE CORPUS

Below section is providing the tabular details of the different content types of the Punjabi Sentence Aligned Speech Corpus. These figures may be helpful in tuning the corpus for various purposes of training, testing and evaluating various algorithms as well as provide useful insights into the dataset. The total duration of Punjabi Sentence Aligned Speech Corpus is 52:24:51 (hh:mm:ss) comprising 31,338 audio segments from 449 speakers.

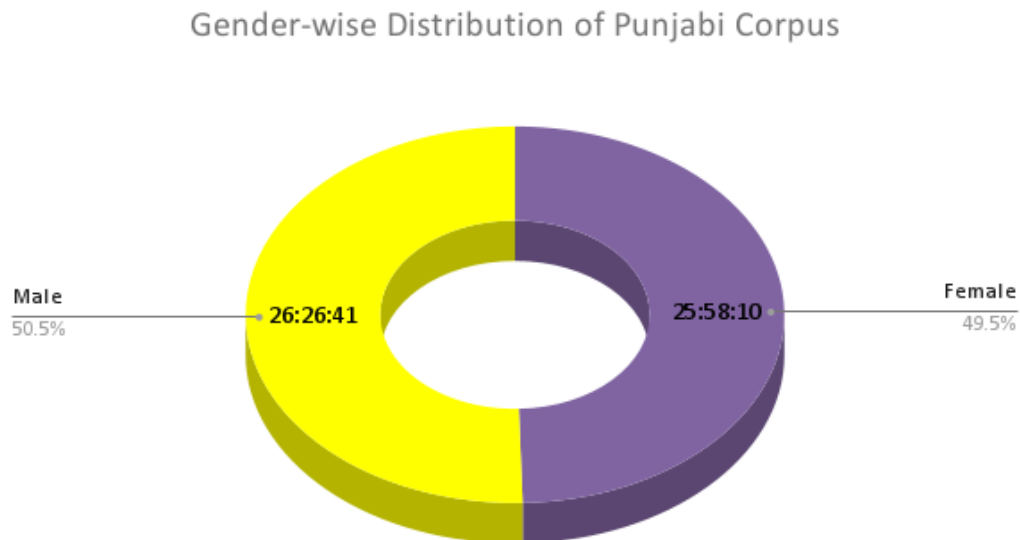


Figure 1: Gender-wise Distribution of Punjabi Corpus

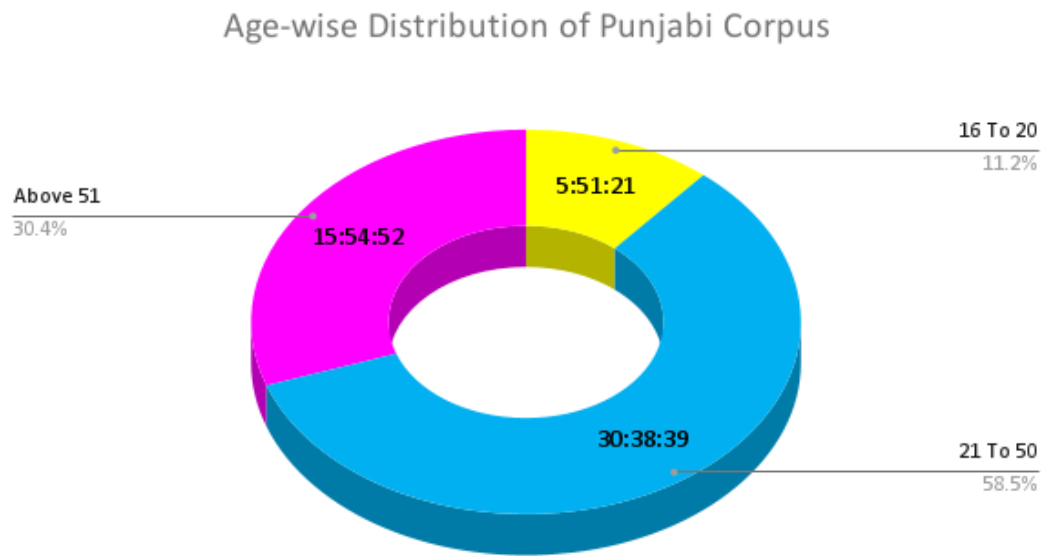


Figure 2: Age-wise Distribution of Punjabi Corpus

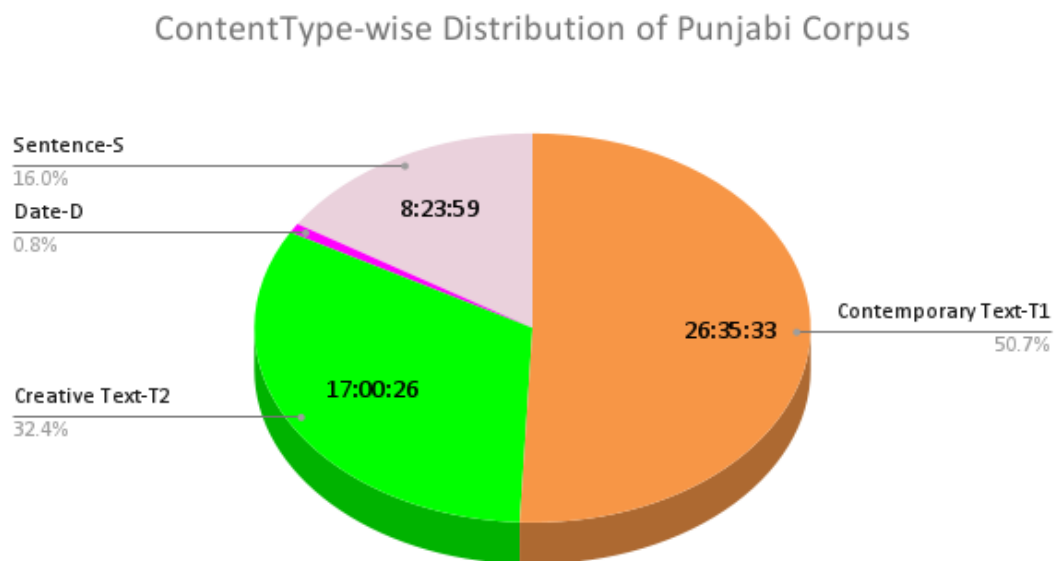


Figure 3: Content Type-wise Distribution of Punjabi Corpus

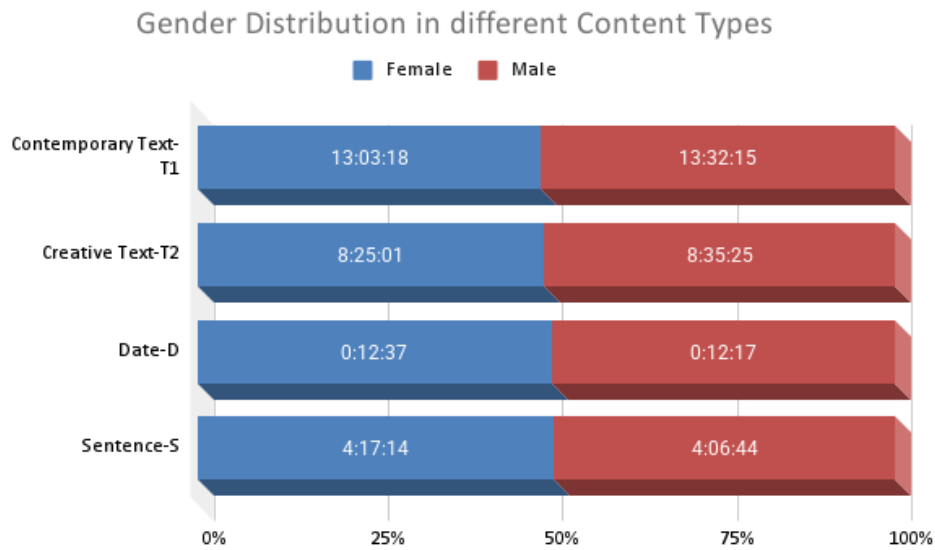


Figure 4: Gender Distribution in different Content Types of Punjabi Corpus

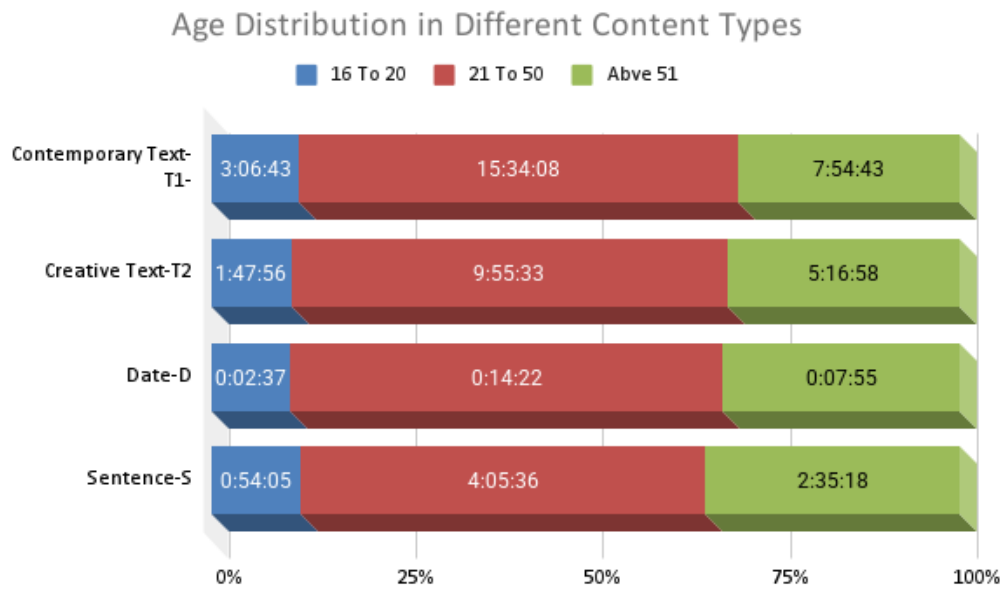


Figure 5: Gender Age Distribution in different Content Types of Punjabi Corpus

### 1.6.1 DURATION OF PUNJABI SENTENCE ALIGNED SPEECH DATA

The table below shows the duration of each of the content types and their distribution across a few factors in Punjabi Sentence Aligned Speech Data.

| Content Type         | Gender | Age Group | Duration (hh:mm:ss.ms) |                 |                 |
|----------------------|--------|-----------|------------------------|-----------------|-----------------|
| Contemporary Text-T1 | Female | 16To20    | 01:44:34.688610        | 13:03:17.679040 | 26:35:32.961558 |
|                      |        | 21To50    | 07:27:18.357540        |                 |                 |
|                      |        | Above51   | 03:51:24.632891        |                 |                 |
|                      | Male   | 16To20    | 01:22:07.587392        | 13:32:15.282518 |                 |
|                      |        | 21To50    | 08:06:49.817238        |                 |                 |
|                      |        | Above51   | 04:03:17.877888        |                 |                 |
| Creative Text-T2     | Female | 16To20    | 01:00:41.990339        | 08:25:00.816673 | 17:00:26.167396 |
|                      |        | 21To50    | 04:54:26.728843        |                 |                 |
|                      |        | Above51   | 02:29:52.097490        |                 |                 |
|                      | Male   | 16To20    | 00:47:13.591862        | 08:35:25.350723 |                 |
|                      |        | 21To50    | 05:01:05.904182        |                 |                 |
|                      |        | Above51   | 02:47:05.854680        |                 |                 |
| Date-D               | Female | 16To20    | 00:01:26.846645        | 00:12:36.534899 | 00:24:53.424891 |
|                      |        | 21To50    | 00:07:20.688374        |                 |                 |
|                      |        | Above51   | 00:03:48.999881        |                 |                 |
|                      | Male   | 16To20    | 00:01:10.203477        | 00:12:16.889991 |                 |
|                      |        | 21To50    | 00:07:00.816734        |                 |                 |
|                      |        | Above51   | 00:04:05.869780        |                 |                 |
| Sentence-S           | Female | 16To20    | 00:29:24.600706        | 04:17:14.171462 | 08:23:58.615162 |
|                      |        | 21To50    | 02:29:00.253382        |                 |                 |
|                      |        | Above51   | 01:18:49.317374        |                 |                 |
|                      | Male   | 16To20    | 00:24:40.033733        | 04:06:44.443699 |                 |
|                      |        | 21To50    | 02:25:36.220200        |                 |                 |
|                      |        | Above51   | 01:16:28.189767        |                 |                 |

Table 1: Representation of Punjabi Sentence Aligned Speech Data Duration

## 1.7 SUMMARY OF SPEAKERS

The table below shows the total number of speakers and their distribution in the Punjabi Sentence Aligned Speech Data.

| Age Group | Female | Male | Total |
|-----------|--------|------|-------|
| 16To20    | 27     | 23   | 50    |
| 21To50    | 133    | 134  | 267   |
| Above51   | 65     | 67   | 132   |
| Total     | 225    | 224  | 449   |

Table 2: Distribution of Speakers of Sentence Aligned Speech Data

## 1.8 REFERENCES

1. Choudhary, N. and D. G. Rao. 2020. The LDC-IL Speech Corpora. In Proceedings of 23rd Conference of the Oriental COCOSDA International Committee for the Coordination and Standardisation of Speech Databases and Assessment Techniques (O-COCOSDA), Yangon, Myanmar, 2020. pp. 28-32, doi: <https://doi.org/10.1109/O-COCOSDA50338.2020.9295011>
2. Choudhary, N. 2021. LDC-IL: The Indian Repository of Resources for Language Technology. *Language Resources & Evaluation*. Springer, Vol. 55, Issue 1. doi: <https://doi.org/10.1007/s10579-020-09523-3>
3. Choudhary, Narayan, Rajesha N., Manasa G. & L. Ramamoorthy. 2019. “LDC-IL Raw Speech Corpora: An Overview” in *Linguistic Resources for AI/NLP in Indian Languages*. Central Institute of Indian Languages, Mysore. pp. 160-174.
4. Ramamoorthy, L., Narayan Choudhary, Poonam Dhillon & Sarbjeet Kaur. 2019. *Punjabi Raw Text Corpus*. Central Institute of Indian Languages, Mysore.
5. Ramamoorthy, L., Narayan Choudhary, Poonam Dhillon, Sarbjeet Kaur & Sandeep Singh. 2019. *A Gold Standard Punjabi Raw Text Corpus*. Central Institute of Indian Languages, Mysore.