

कोंकणी – वाक्य संरेखीत वाक कॉर्पस



**KONKANI – SENTENCE ALIGNED
SPEECH CORPUS**

KONKANI 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

Title: Konkani Sentence Aligned Speech Corpus

Authors: Saurabh Varik, Rajesha N., Manasa G., Srikanth D., Stephen Fernandes, Nithin S.,
Narayan Kumar Choudhary, Shailendra Mohan

ISBN: 978-81-19411-62-7

CIIL Publication No.: 1429

First published: AD 2023 September
Bhadrapada 1945 Śaka

© Central Institute of Indian Languages, Mysuru 2023

Publisher: Prof. Shailendra Mohan, Director, CIIL

Production Team

Head, Publication Unit: Umarani Pappuswamy

Officer-in-Charge, Publication Unit: Aleendra Brahma

Artist: H. Manohara

Staff-in-charge: R. Nandeesh

Compositor: M. N. Chandrashekar

Cover design: Saurabh Varik

TABLE OF CONTENTS

Table of Contents.....	iv
Figures.....	iv
Tables.....	iv
1 Speech Annotation.....	5
2 Konkani Speech Annotation.....	9

FIGURES

Figure 1: Gender-wise Distribution of Konkani Corpus.....	12
Figure 2: Age-wise Distribution of Konkani Corpus.....	12
Figure 3: Content Type-wise Distribution of Konkani Corpus	12
Figure 4: Gender Distribution in different Content Types of Konkani Corpus.....	13
Figure 5: Age Distribution in different Content Types of Konkani Corpus	13

Tables

Table 1: Representation of Konkani Sentence Aligned Speech Data Duration.....	14
Table 2: Distribution of Speakers of Konkani Sentence Aligned Speech Data.....	14

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 ‘avanviittipoyi’ ‘He went home’ in an informal way then it is annotated as ‘avanviittilpoyi’ 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 “kekeralam”. 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 KONKANI SPEECH ANNOTATION

Saurabh Varik, Narayan Kumar Choudhary

2.1 OVERVIEW OF SENTENCE ALIGNED SPEECH CORPUS

Konkani Sentence Aligned Speech Corpus is created by annotating the speech data collected by LDC-IL. A detailed explanation of the Konkani Raw Speech Corpus will be available in the [Konkani Speech Data Documentation](#) (Ramamurthy, L. et. Al, 2019). LDC-IL Konkani Sentence Aligned Speech files contains an audio file and two textual layers in Konkani 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

‘Konkani_Female_16To20_Contemporary Text-T1_SP-0031_T1-0031-001.wav’

LDC-IL Sentence Aligned Speech corpus for Konkani 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 list - each speaker has typically recorded 25 sentences randomly selected from this set. Date format contains a question uttered by the investigator and the response of 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.

2.2 OBSERVATIONS

LDC-IL sentence-level speech annotation strictly follows what the speaker pronounces to produce the phonetically normalised annotation. The text has been written in the official script of the language and the speech is transcribed as narrowly as the script supports. Even if it is read speech data, there are widespread variations in the pronunciation. For example, speakers from different regions speak the same word in different ways. For example, *uṅṅ kər* "reduced" sometimes sounds like (*uṅṅ ŋkər*), Kokum (*Garcinia indica*) - *bʰɪrɪŋd̪ə/bʰɪrəŋd̪ə/bʰɪŋd̪ə*.

There were also variations in how numbers were pronounced. For example, while reading sports news, the speakers read scores of different sports such as cricket, tennis, etc in different ways and they deviated from the standardised way of pronouncing the scores. Similarly, there were some errors in reading large numbers such as thousands or lakhs and also in reading decimals, fractions, etc. Most of the speakers faced difficulty in pronouncing foreign names which frequently appear in sports news. Abbreviations and rarely-used words also influenced the reader’s fluency.

2.2.1 PHONETIC ALTERNATION IN KONKANI SPEECH DATA

Read speech has disfluencies like unwanted pauses, elongated syllables, word fragments, self-corrections, and repeated words. Some such disfluencies in the recording are given below:

a. Repetition of words

While reading, if the informant observes that the word has been pronounced not in correct or effective manner then normally the speaker repeats part of the word, whole word or the phrase. Sometimes the speaker was struggling to read the text and repeats when the content is about unfamiliar subjects or there are many foreign words or words which are difficult to pronounce. These are mainly instances of self-correction.

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.: g^hə-g^hərə, gə-g^hərə

c. Intended speech

Intended speech occurs when the speaker slows down or fastens up their speech. Typically, it happens at the end of the sentence. It has resulted in inaudible speech in some instances. For example, if the audio is transcribed as mət̪rɑ:[ləjɑ:]*t̪e:rə (मंत्राचेर*[ल्या]), it shows that [ləjɑ:] (ल्या) is not properly audible. In some longer words, the middle of the syllable or phone might not be audible to the listener or are skipped by the speaker. For example, in prəḍ^hɑ:ṇə*ṭri:[प्रधान*त्री] the middle part is not audible.

d. Addition and Deletion

An extra vowel or a consonant or a syllable is sometimes added into a word. The sound which already exists in the word might be repeated or a different sound might be inserted into the word. E.g.: s̪ət̪ɑ:rə (संचार) > s̪ət̪ət̪ɑ:rə (संचचार); s̪unɑ:pərɑ:ṭə (सुनापरांत) > s̪unɑ:pərɑ:ṇəṭə (सुनापरांतत)

Deletion or elision of a vowel or a consonant or a syllable from a word is also a common phenomenon attested in the corpus.

E.g.: kəru:ṇə [करून] > kərəṇṇə [करन]

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.: ḍille:>me:[dille>मेळळे]

Dissimilation is dropping out a syllable or a letter by the influence of adjacent speech segments.

E.g.: p^hərəmɑ:ṇə > pərəmɑ:ṇə फरमान>परमान

f. Colloquial usage

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

E.g.: a:t̪^hi:ṣə > əḍəṭi:ṣə आट्टीस>अडतीस, ikəra: > əkəra: इकरा > अकरा

g. Lengthening and Shortening

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

E.g.: a:ɲi: > a:ɲi आनी > आनि, kəvɪt̪a: > kəvi:t̪a: कविता > कवीता

h. Substandard alternation

It has been observed that some speakers have consistently | replaced the aspirated sounds with their unaspirated counterparts. Moreover, in Kozhikode region people adapted easily articulated sounds.

E.g.: pa:kʰo: > pa:ko:पाखो > पाको, t̪a:ri:kʰə > t̪a:ri:kə तारीख > तारीक

i. Final vowel Elision

In the Kanara-Saraswat dialect, all solitary words end with a vowel, but in connected speech all word-final vowels in multisyllabic words are omitted when another word follows without a pause. hãva tãkkã āpaytã (hãva "me", tãkka "he", āpaytã "telephone") is pronounced hãv tãk āpaytã. Such vowel omissions in connected speech are also found in the Sashti dialect of Christianity. In both dialects, if the omitted vowel is a prevowel, the preceding consonant is palatalized.

Vowel Rounding in Christian Dialects: In the dialects of Christianity (Valdez Christian and Saxti Christian), the vowel a is fused with the o, resulting in fewer vowel phonemes.

j. Compound word splitting

Long agglutinated words have been read in such a way that a pause is at the point of joining and that interrupts the natural flow of language.

E.g.: ba:jələməɲi:ʃə > ba:jələ məɲi:ʃə, gʰərəɖa:rə > gʰərə ɖa:rə

2.3 SUMMARY OF THE CORPUS

The total duration of Konkani Sentence Aligned Speech Corpus is 83:19:42 (hh:mm:ss) comprising 34,091 audio segments from 487 speakers. Figures 1, 2 and 3 show the distribution of the corpus with respect to gender, age and content type, respectively. Figures 4 and 5 show gender and age distributions for each content type respectively. Table 7 gives a break-up of the corpus in terms of recordings obtained from different kinds of texts and also other demographic details. Table 8 shows the age and gender-wise distribution of all the speakers.

Gender-wise Distribution of Konkani Corpus

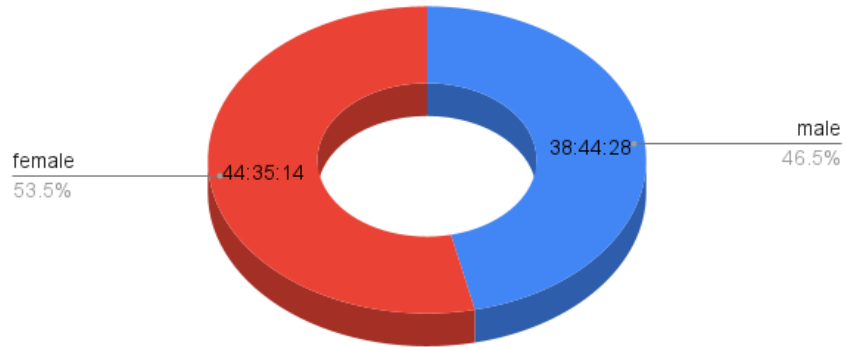


Figure 1: Gender-wise Distribution of Konkani Corpus

Age-wise Distribution of Konkani Corpus

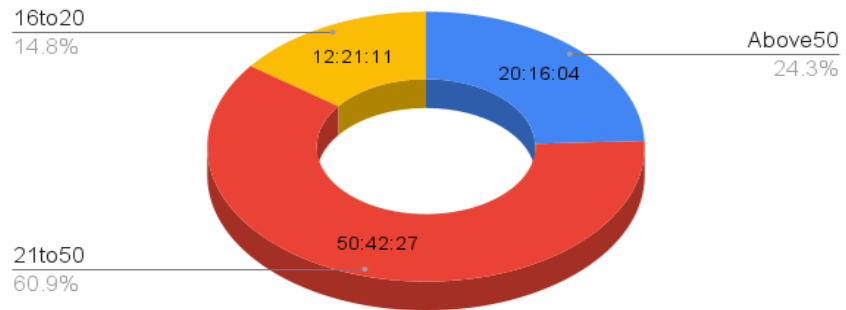


Figure 2: Age-wise Distribution of Konkani Corpus

Content Type-wise Distribution of Konkani Corpus

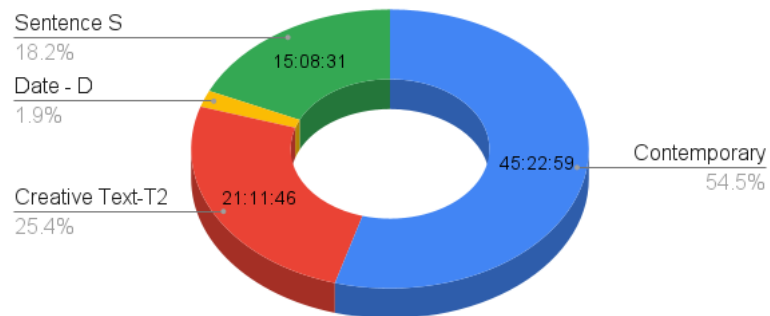


Figure 3: Content Type-wise Distribution of Konkani Corpus

Gender Distribution in different Content Types

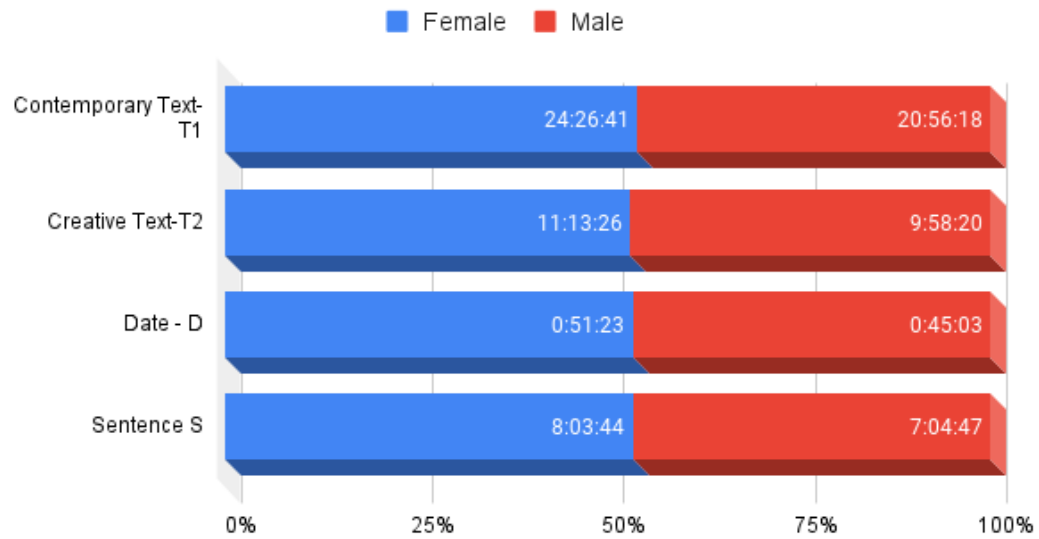


Figure 4: Gender Distribution in different Content Types of Konkani Corpus

Age Distribution in different Content Types

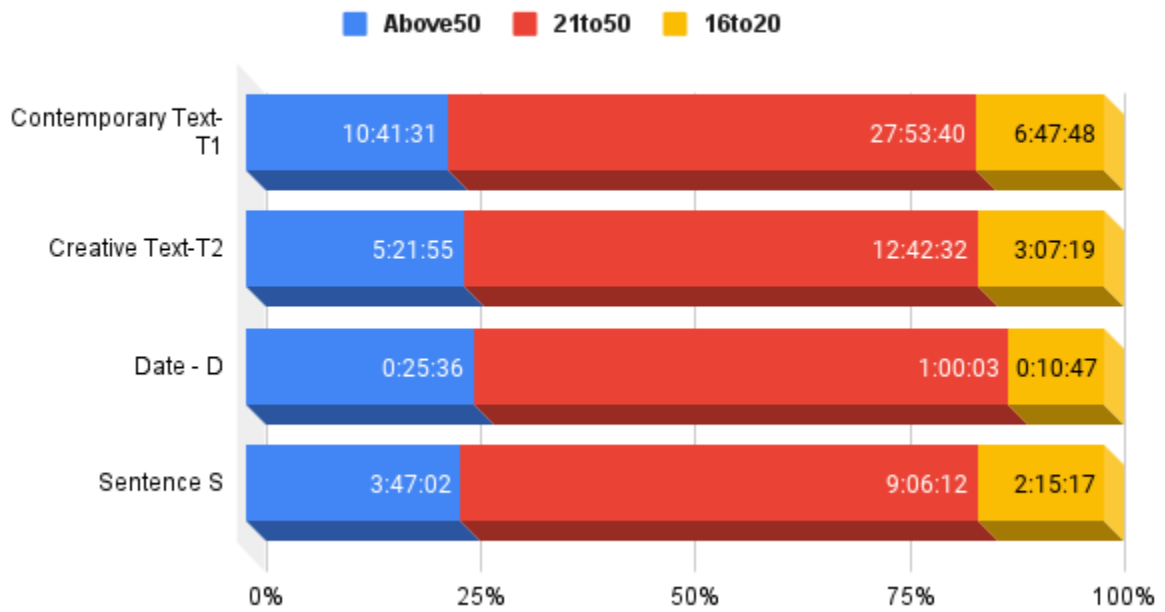


Figure 5: Age Distribution in different Content Types of Konkani Corpus

2.3.1 DURATION OF KONKANI SENTENCE ALIGNED SPEECH DATA

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

Content Type	Gender	Age Group	Duration (hh:mm:ss.ms)		
Contemporary Text-T1	Female	16To20	03:56:35.127344	24:26:41.071740	45:22:58.516021
		21To50	15:07:25.134619		
		Above51	05:22:40.809776		
	Male	16To20	02:51:12.656557	20:56:17.444281	
		21To50	12:46:15.072803		
		Above51	05:18:49.714921		
Creative Text-T2	Female	16To20	01:49:03.351690	11:13:26.284276	21:11:46.954730
		21To50	06:39:52.328808		
		Above51	02:44:30.603779		
	Male	16To20	01:18:15.916249	09:58:20.670454	
		21To50	06:02:40.489887		
		Above51	02:37:24.264317		
Date-D	Female	16To20	00:06:51.936300	00:51:22.595258	01:36:24.983446
		21To50	00:31:17.980378		
		Above51	00:13:12.678580		
	Male	16To20	00:03:54.638059	00:45:02.388189	
		21To50	00:28:45.135341		
		Above51	00:12:22.614789		
Sentence-S	Female	16To20	01:20:13.004817	08:03:44.234407	15:08:31.069589
		21To50	04:48:10.121925		
		Above51	01:55:21.107665		
	Male	16To20	00:55:03.584759	07:04:46.835182	
		21To50	04:18:01.760069		
		Above51	01:51:41.490353		

Table 1: Representation of Konkani Sentence Aligned Speech Data Duration

2.4 SUMMARY OF SPEAKERS

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

Age Group	Female	Male	Total
16To20	40	26	66
21To50	157	140	297
Above51	62	62	124
Total	259	228	487

Table 2: Distribution of Speakers of Konkani Sentence Aligned Speech Data

2.5 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 Co-ordination 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, Saurabh Varik, Rashmi Shet Tanawade & Yashwant D Gawas. 2019. A Gold Standard Konkani Text Corpus. Central Institute of Indian Languages, Mysore.
5. Ramamoorthy, L., Narayan Choudhary, Saurabh Varik & Rashmi Shet Tanawade. 2019. Konkani Raw Speech Corpus. Central Institute of Indian Languages, Mysore.