

খুং চংদ্রিবা মনিপুৰী খোঙেলকোন

Manipuri Raw Speech Corpus



MANIPURI RAW SPEECH CORPUS

Manipuri Raw Speech Corpus

Authors:

Mr. Amom Nandaraj Meetei
Ms. Yumnam Premila Chanu
Mr. Rajesha N.
Ms. Manasa G.
Dr. Narayan Choudhary
Dr. L. Ramamoorthy



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
Mysore, India-570006

Manipuri Raw Speech Corpus

This Documentation is a part of LDC-IL Manipuri Raw Speech Corpus.

For Further contacts: oic@ldcil.org

All rights reserved

© Central Institute of Indian Languages, Mysore-2018

Authors:

Mr. Amom Nandaraj Meetei

Ms. Yumnam Premila Chanu

Mr. Rajesha N.

Ms. Manasa G.

Dr. Narayan Choudhary

Dr. L. Ramamoorthy

ISBN: 978-81-7343-247-7 Electronic publication of Corpus

Last updated: December 12, 2018

Publisher:

Linguistics Data Consortium for Indian Languages

Central Institute of Indian Languages

Mysore, India-570006

Contents

1	LDC-IL Raw Speech Corpora: An Overview	1
1.1	Introduction	1
1.2	LDC-IL Speech Corpus	2
2	Content Type Descriptions.....	3
2.1	T1: Contemporary Text	3
2.2	T2: Creative Text	4
2.3	D: Date	4
2.4	S: Sentences	4
2.5	W1: Command and Control Words.....	5
2.6	W2: Proper Noun (Person Names and Place Names)	5
2.6.1	Person Names	5
2.6.2	Place Names.....	5
2.7	W3: Most Frequent Words	6
2.7.1	W3A: Most Frequent Words-Part	6
2.7.2	W3B: Most Frequent Words-Full	6
2.8	W4: Phonetically Balanced Vocabulary	6
2.9	W5: Form and Function words	7
3	Planning for Fieldwork	8
3.1	Dataset preparation and distribution	8
4	Field Work	10
4.1	Possible places for collecting data	10
4.2	Field work Ethics	10
4.3	Data Collection.....	11
4.3.1	Technical Specifications for collecting data	11
4.3.2	Metadata.....	12
4.3.3	Data Transferring and Storing.....	15
4.3.4	Observations	15
5	Organising and Archiving the Data	16
5.1.1	Text - Speech Mapping and Naming Conventions	16
5.1.2	Observations	16

6	Data verification and Quality Control	17
7	Manipuri Raw Speech Corpus	18
7.1	Introduction	18
7.1.1	Manipuri in the Eighth Schedule of the Constitution of India	18
7.1.2	Area and Population of Manipur	18
7.1.3	A cursory Glance at Manipuri Dialects	19
7.1.4	Intra-regional speech variation and LDC-IL dataset.....	19
8	Dataset preparation for Manipuri.....	21
9	Transliterations in LDC-IL Manipuri Read corpus.....	22
10	Summary of the Corpus	23
10.1	Summary of the Audio Segments	23
10.2	Duration of the Raw Speech Data.....	24
10.3	Distinct Set	24
10.3.1	The Contemporary Text (News)- T1.....	24
10.4	Random Set.....	25
10.4.1	The Created Text-T2.....	25
10.4.2	The Sentences-S.....	25
10.4.3	The Date-D	25
10.4.4	Command and Control Words-W1	26
10.4.5	Person Names –W2.....	26
10.4.6	Place Names-W2	26
10.4.7	Most Frequent Words-PART-W3A	27
10.5	Full Set.....	27
10.5.1	Most Frequent Words-Full-W3B.....	27
10.5.2	The Phonetically Balanced Words-W4.....	27
10.5.3	The Form and Function Words-W5.....	28
10.6	Native Speakers Distributions.....	28
10.7	Conclusion.....	28
10.8	Reference.....	28

Table of Figures

Table 1: LDC-IL Speech Data Content Types	3
Table 2: Standard Speech Dataset Distribution for Each LDC-IL Fieldwork with Modle-1 Dataset	8
Table 3: Standard Speech Dataset Distribution for Each LDC-IL Fieldwork with Modle-2 Dataset	9
Table 4: Technical Specifications for collecting data	11
Table 5: Metadata Legends and their Description.....	14
Table 6: LDC-IL Speech Dataset.....	21
Table 7: Table of Contents in LDC-IL Dataset.....	21
Table 8: Four Phases of Speech Data Collection	21
Table 9 : Manipuri Audio Segments and their Distribution.....	23
Table 10: Duration of the Collected Manipuri Speech Data.....	24
Table 11: Distribution of Manipuri Contemporary Text (News) Data	24
Table 12: Distribution of Manipuri Created Text:T2.....	25
Table 13: Distribution of Manipuri Sentences	25
Table 14: Distribution of Manipuri Date Format	25
Table 15: Distribution of Manipuri Command and Control Words.....	26
Table 16: Distribution of Maniprui Personal Names	26
Table 17: Distribution of Manipuri Place Names	26
Table 18: Distribution of Maipuri Most Frequent Words – Part	27
Table 19: Distribution of Manipuri Most Frequent Word-Full.....	27
Table 20: Distribution of Manipuri Phonetically Balanced Words-W4	27
Table 21: Distribution of Manipuri Form and Function Words-W5.....	28
Table 22: Distribution of Manipuri Native Speakers.....	28

1 LDC-IL RAW SPEECH CORPORA: AN OVERVIEW

1.1 INTRODUCTION

Lack of basic linguistic resources have been the first and foremost bottleneck in development of language technology for Indian languages. When text data itself has been available for most of the Indian languages, one could not even think of the speech data. India is one of the foremost multilingual country where multilingualism is ingrained and most people speak more than one language with more than 75 languages having more than one million speakers (as per 2011 Census of India data). As per a study¹ of KPMG and Google released in 2017, the internet user base grew at a compound annual growth rate (CAGR) of 41% between 2011 and 2016 to reach 234 million users at the end of 2016 and this trend is likely continue. It is also estimated that internet users in Indian language will account for nearly 75% of India's internet user base by 2021.

Despite this, the availability of technology in Indian languages have been on close to null. This is mainly due to the reason that the technology developing agencies find it either too difficult to come up with the language support on various applications for Indian languages or it is economically not a viable solution. However, recent analyses from various quarters have shown that the latter is not correct and the major issue is availability of the linguistic resources based on which language technology and language support for various types of applications proves to be a bottleneck for the developing community, be it industry or otherwise.

Considering this as an issue, the Government of India has taken several initiatives to provide the basic ingredients which may prove as a catalyst for the development of language technology in Indian languages. As part of the this initiative, a scheme named Linguistic Data Consortium for Indian Languages (LDC-IL) was established by the Ministry of Human Resource and Development at Central Institute of Indian Languages, Mysore.

The goal of LDC-IL was to develop linguistic resources for all Indian languages with the initial focus more on the scheduled languages of India. These linguistic resources may be as deemed fit by the language technology developing community.

Based upon the recommendations of the Project Advisory Committee which includes ex-officio members from MeitY, IITs Ministry of HRD, Director and other academicians from reputed Institutes/Universities working in this area as well as major and minor industrial entities working in this area, the LDC-IL decided to embark upon developing the text and speech corpus for the scheduled languages of India.

There have been several types of datasets prepared under LDC-IL. This document serves as a generic documentation for the raw speech corpus of the LDC-IL being released for several languages.

¹<https://assets.kpmg.com/content/dam/kpmg/in/pdf/2017/04/Indian-languages-Defining-Indias-Internet.pdf>

1.2 LDC-IL SPEECH CORPUS

LDC-IL speech corpus is collected after careful deliberations on what type of speech corpus is required for various types of speech based linguistic analysis that may suit multifarious needs of the research and development community.

After several meetings with the experts from around India and abroad, it was decided that LDC-IL should focus on not just developing a speech corpus for a particular need, rather to get the data that would be useful for various tasks such as ASR, STT, linguistic analysis, speech therapy and so on.

Keeping this in mind, various types of content were created a priori before the speech recordings took place. The content of these datasets have been prepared in consultation with the experts from the language as well as linguists giving inputs to ensure that no specific sound patterns are missed out.

For example, it has been ensured that the speech datasets contain all the phones and allophones of the language and ample examples are available in the language to prove their phonemic status in the language. To ensure that the corpus is good for an ASR, the continuous speech is recorded in natural environment.

2 CONTENT TYPE DESCRIPTIONS

Each content type has a number of files with each file containing standard content. A sub-set of these files in each of the content types selected randomly constitute subsets that are given to a speaker for reading out in natural flow. A few full sets (namely W3B, W4, and W5) are also read in full by certain selected speakers in each age group.

There are three age group ranges selected for LDC-IL datasets. These are ‘16 to 20 years’, ‘21 to 50 years’ and ‘above 50 years’. Attempt has been made to collect equal number of male and female data from each of the age groups.

The list of the datasets and their notation is given in the table below:

SL	Notation	Content Type
01.	T1	Contemporary Text (News)
02.	T2	Creative Text
03.	S	Sentence
04.	D	Date
05.	W1	Command and Control Words
06.	W2	Place Name
07.	W2	Person Name
08.	W3A	Most Frequent Word-Part
09.	W3B	Most Frequent Word-FullSet
10.	W4	Phonetically Balanced-Fullset
11.	W5	Form and Function Word-Fullset

Table 1: LDC-IL Speech Data Content Types

Detailed descriptions of each of the content types are given in the following sub-sections.

2.1 T1: CONTEMPORARY TEXT

The Contemporary Text (news) data is given the notation of T1. News items have been selected from the LDC-IL news corpora. The text is contemporary in nature as the news items such have been picked over a period from 2005 to 2012, either from news websites or from the print editions newspapers of the respective language.

The domain information is present in the news items as well as the news items deal various topics such as political news, editorials, sports news and so on. Given that the news items have been collected from local news reported for each language, the style may be considered as colloquial or belonging to the news reporting style.

Each LDC-IL dataset ‘Contemporary Text’ contains minimum of 500 words per speaker, which is rarely repeated. Since it is the continuous text, it constitutes the largest part of the speech corpora, in terms of data size and time duration.

2.2 T2: CREATIVE TEXT

‘Creative Text –T2’ is extracted mainly from literary sources. It is used to capture literary terms. Creative Texts are stories or essays collected from books. The text may be any standard text which is descriptive in nature. It exhibits the language style of a particular period from which the text is taken.

Creative texts were prepared in two types. In the first 6 or 8 essays or short stories were prepared. One of these selected randomly from the set, is assigned to one speaker for reading out. The same story may be read by multiple speakers.

In the other approach a distinct text is given to each individual

The creative text section of the LDC-IL Speech dataset comprises of mostly six essays or short stories. One of these essay or short story, selected randomly from the set of the six stories, is assigned to one speaker for reading out. The same story may be read out by multiple speakers.

2.3 D: DATE

Languages tend to speak out the date in a specific and many a times in a particular manner which may not always conform to the grammatical structure of the language. To capture it, LDC-IL tried to document how a date is spoken in each of the languages.

The normal way is put a question before the informant the answer of which must be in a date format. Normally the following six questions were placed before the informant and the informants would answer minimum one of the questions.

1. What is tomorrow’s date?
2. When is Gandhi Jayanthi observed?
3. What is the date today?
4. When do we celebrate our Independence Day?
5. What is your date of birth?
6. On which date you go to market?

2.4 S: SENTENCES

To ensure that all the types of syntactic structures are covered in the speech data, a set of sentences have been constructed with the help of language experts and linguists for each of the languages. It is ensured that all possible sentence structures are covered including all types of tenses, aspects, moods, compound and complex sentences and so on.

These sentences are isolated sentences and not part of a continuous speech. While care has been taken to extract sentences from the text corpus of the corresponding language, sometimes sentences have also been modified to ensure that the specific valid sentence structure of the language is present.

Very long sentences are avoided while selecting or constructing the sentences, so that the informant can read the sentences easily. The words used in these sentences are common words which are found in day-to-day life. Each sentence in the list contains minimum four words. The sentences are not too long so that each sentence does not span for more than a line in the prompting sheet. Care is taken to avoid abusive or taboo words.

Each speaker is given 25 sentences out of this sentence list for reading out.

2.5 W1: COMMAND AND CONTROL WORDS

Spoken language usually contains a lot of sentences that are commands or use a lot of control words. This happens mostly in the conversational speech. Even though the LDC-IL speech corpus at present does not contain the conversation speech, an attempt has been made by including common command and control sentences/phrases carefully crafted with the help of respective language experts and linguists.

These include imperative sentences, optative sentences as well as other controlling phrases which may come as a reply to an interrogative sentence. Each of the languages has a set of command and control sentences created before the speech data is recorded. Each speaker is given a list of 30 command and control sentences randomly selected from the set. Each of these phrases/sentences is repeated three times by each speaker while recording.

2.6 W2: PROPER NOUN (PERSON NAMES AND PLACE NAMES)

Recognizing proper nouns by using an ASR system is a complex task. For example, voice recognition application in mobile phone may have a few hundreds of names to distinguish when placing a call through voice command. Native speakers use different pronunciations depending on their language of origin and familiarity with the language. The speakers use different pronunciation for native and foreign names ranging from a nativised pronunciation to a totally foreignised pronunciation. All this adds to the complexity of an ASR system in recognizing proper nouns. To address this issue LDC-IL speech data has been collected to include person names and place names.

2.6.1 Person Names

Person names are included to capture the native pronunciations. The names are taken from people from different walks of life like Politicians, Film Actors and Directors, Writers, Kings and Queens, Astrologers, Historical Personalities, Scientists, Sports persons etc.

2.6.2 Place Names

Place names are included to capture the native pronunciations. This data set contains Indian place names. These include main cities, district names and popular tourist destinations from all over India. Some local place names are also included like names of villages, taluk headquarters, district names, local forest reserves, local tourist and pilgrimage destinations etc.

Each speaker typically pronounce 20 person names and 10 place names, out of the total Proper Noun wordlist of the particular language. Each word is uttered three times

2.7 W3: MOST FREQUENT WORDS

Most frequent word list is the regularly and repeatedly used list of words. Since these words are used most frequently in a language, it is imperative to have these words in our dataset.

The most frequent words dataset is derived from LDC-IL Corpus. However, it may be noted that when the most frequent word list was extracted, the text corpus was rather small. So, the frequency list might change if it is compared to the current LDC-IL text corpus.

2.7.1 W3A: Most Frequent Words-Part

The most frequent words of a language are randomly picked from a list of around 1000 most frequent wordlist of a language. Each speaker records randomly selected 30 words from this list. Each word is uttered thrice.

2.7.2 W3B: Most Frequent Words-Full

Two speakers, one male and one female, pronounces the full set of 1000 most frequent words. This is done for each dialect of the language, if available.

2.8 W4: PHONETICALLY BALANCED VOCABULARY

To cover all possible phonemic occurrences of a language, the “phonetically balanced Vocabulary” is prepared. It is a list of words in which the occurrence of a phoneme in initial medial and final positions of that language can be represented.

The pronunciation of the phoneme is varied according to the position of the phoneme in a word and the influence of the following and proceeding phoneme. The pronunciation of initial position is different from middle and final positions. For example the phoneme ‘pa’ is used in different forms while pronouncing words like

- ‘**pallavi**’- ‘pa’ inherent vowel at initial position (CV initial)
- ‘**prakata**’ - ‘p’ as pure consonant in conjunction with ‘ra’ in initial position, (CCV Initial)
- ‘**spandana**’, - ‘pa’ with inherent vowel preceded by a consonant at medial position (CCV Initial)
- ‘**parikalpane**’- ‘pa’ inherent vowel at initial position (CV initial) and ‘pa’ with inherent vowel preceded by a consonant in the medial position (CCV Medial)
- ‘a:**pta**’ - ‘p’ with followed by a consonant in the final position (CCV medial)

Using the articulation score as the measure, phonetically balanced lists have been used to test differences among transmission systems and to test the effects of noise. The phonetically balanced words used in word recognition testing contain speech sounds that occur in the same frequency as those of conversational speech.

2.9 W5: FORM AND FUNCTION WORDS

Form and function words dataset is a closed class list of words. They are quite limited in a language. These constitute mostly the indeclinable words of a language. Form words are static, bearing some content with them. They are meaningful and are actually the building blocks of a language.

The Form and Function dataset includes Grammatical function words, numerals, kinship terms, measurement terms, list of colours, days, months, seasons, directions, zodiac signs, body parts, planets etc. These words are included to the native words which may not be frequent in the overall corpus, but needs representation.

3 PLANNING FOR FIELDWORK

3.1 DATASET PREPARATION AND DISTRIBUTION

To ensure representativeness of the speech corpora, a conscious effort has been made to balance the speech data by taking varieties of styles into consideration. The first and foremost among at LDC-IL has been to take an expert view on the varieties of languages. For example, for Kannada it is ensured that speech varieties from different regions such as Hyderabad Karnataka, Bombay Karnataka, Coastal Karnataka and Old Mysore get proportionate weightage.

LDC-IL collected the data using two approaches, with two different kind of Dataset Models They are as follows

- Dataset Model 1 (T1, T2, W1, W2, W3, W4, W5, S, D)
- Dataset Model 2 (Distinct Texts of T1 and T2)

Some Languages followed Model-1 only, and some Languages followed both Model-1 and Model-2 After the regions are identified, speech samples are collected as per the criteria shown in the tables below:

Standard Speech Dataset Distribution for Each LDC-IL Fieldwork Dataset Model 1							
Content type	Content size#	Content to be read by one speaker	Total No. of speakers	Age group wise no. of speaker; Female & Male equally distributed#			Content selection type
				16-20	20-50	50+	
Contemporary Text	150 Texts	1 Text	150	18	90	42	Distinct Text
Creative Text	6 Texts	1 text	150	18	90	42	Random set*
Sentences	142 Sentences	25 Sentences	150	18	90	42	Random set*
Command and Control Words	82 Words	30 Words	150	18	90	42	Random set*
Person Names	489Words	20 Words	150	18	90	42	Random set*
Place Names	511 Words	10 Words	150	18	90	42	Random set*
Most Frequent Words	1144Words	30 Words	150	18	90	42	Random set*
Phonetically Balanced Vocabulary	390 words	Full set	6	2	2	2	Full set to be read by the speaker
Form and Function Words	432 words	Full set	6	2	2	2	Full set to be read by the speaker
1000 Most Frequent Words	1000 Words	Full set	2	0	2	0	Full set to be read by the speaker
*picked randomly by machine #The figures shown are for illustration purpose only. The numbers may differ for each language. Please refer specific Language documentation for actual figures.							

Table 2: Standard Speech Dataset Distribution for Each LDC-IL Fieldwork with Modle-1 Dataset

Speech dataset distribution for fieldwork Dataset Model 2						
Content type	Content size	Content to be read by one speaker	Total No. of speakers	Age group wise no. of speaker; Female & Male equally distributed		Content selection type
				16-20	21-50	
Contemporary Text (News)	150 Texts	1 Text	150	75	75	Distinct Text
Creative Text	150 Texts	1 text	150	75	75	Distinct Text

Table 3: Standard Speech Dataset Distribution for Each LDC-IL Fieldwork with Modle-2 Dataset

As the data is collected from different cities across India (as per the demand of the language), its imperative that proper preparation is made before proceeding towards the field such that day-to-day necessities of field are met with. Investigators had to make that s/he had sufficient charged batteries as well as alkaline batteries if so required, empty SD cards, laptops in proper condition, sufficient number of random and full datasets (prompt sheets) and so on. These formed as the daily routine for the linguists while in the field.

4 FIELD WORK

Some common guidelines and instructions were provided to the field workers before proceeding to the field. A brief of it is noted below.

4.1 POSSIBLE PLACES FOR COLLECTING DATA

Once the dataset is prepared and taken to the field, the next step is to determine places where there is an availability of speakers who can read fluently. The best possible places are schools, colleges, universities, govt. offices etc.

The Head of these organizations have to be briefed and asked permission for recording data from students, faculties etc. Certain infrastructural requirements like space, if possible power source for charging batteries etc. has to be requested from the institutions. The speakers from whom we collect data are referred as informants.

4.2 FIELD WORK ETHICS

The informants are briefed about the procedures, nature and purpose of speech data collection. Informants are informed about the funding agency behind the data collection. In case of LDC-IL, the data collection is funded by Govt. of India. Informant are made aware of who exactly is carrying out the data collection process and what will be done with the data collected before they give their consent.

There have been situations where the informant would find it distressing that the data given by them will be segmented and further processed. In such cases, their opinions have to be respected and the investigators have to refrain from taking their data. The informants are made aware of the degree of confidentiality and anonymity that will be maintained after collecting the data. The informants are also made aware of the potential benefits of the data to the wider community. Once the informant is aware of all these information and is ready to give the data, consent is acquired in written along with certain personal details such as their educational qualification, mother tongue, place of elementary education etc.

Informants are allowed to read the dataset earlier before recording so that they can get familiar with the content of the text. It is ensured that the informants do not have any objection to the content they are about to read. For example, the informant may have objection regarding the political, social views expressed in the content. In such cases, a different dataset is offered to the informant. There are certain texts in the data set, which may pose difficulty for a certain informant to read. For Example, some informants may have difficulty in reading contents which involve dialogues between people. Such contents may differ in dialects spoken by the informant which may pose a difficult situation for them while reading. In such cases, a different dataset is offered to the informant. Complex datasets are given only to the informants who are capable of reading them and state likewise.

An attempt is made to reduce the extra noise as much as possible before recording. If necessary, test recordings are conducted before the actual recording on certain portions of the text.

Brief introduction about the informant and investigator along with details like place, time, region etc. are collected at the beginning of each recording. The conversation between investigator and informant is done in their native language so that the informant is comfortable and the natural flow of language is established.

Care is taken while recording the words, so that there is a pause between two words or between utterances of the same word. All the words of the content type W1 to W5 (i.e. ‘Command and Control words’, ‘Proper Nouns’, ‘Most Frequent Words’, ‘Phonetically Balanced Vocabulary’ and ‘Form and Function words’) are repeated three times in a sequence. A pause is maintained between two sentences as well while recording.

While recording the News Item and Creative Text, the informants are briefed to read the text given, as naturally as possible. It should be as natural as reading a book or newspaper. Informants answer to a particular question themselves regarding date format. This is done to capture how people usually pronounce the date. The investigator does not prompt any particular format

4.3 DATA COLLECTION

The LDC-IL data is recorded using Roland EDIROL Recorder. It is a 24-bit Linear PCM (R-09) Recorder.

4.3.1 Technical Specifications for collecting data

Recording Setup:	Sample Rate : 48.0 KH
Recording Mode:	wav -16bit
Date Setup:	Current date and time.
Storage:	SD Card
Power:	• Always use rechargeable batteries (Ni-MH) for recording. Otherwise line hum will come. Never use Ni-CD battery type as it is potential for ‘memory effect’. • Rechargeable batteries need to be thoroughly recharged before recording (minimum 16 hrs continuous charging).
Peak	While recording please be aware that it should not reach the peak i.e. PEAK (in the recorder) should not glow.
Recording Distance	• Keep minimum 5 cm to 25 cm distance between the microphone and the speaker and if possible use microphone holder. • The recorder should not be placed orthogonally but it should be placed diagonally. • Do not move the recorder during recording • Fix the recorder upon a table/ fixed plane if possible. • Try to have fixed the distance between the recorder and speaker • The recorder should not be placed orthogonally but it should be placed diagonally

Table 4: Technical Specifications for collecting data

After each recording, it is recommended to verify the recorded data, whether it is recorded in the right way. If the informant also wishes to hear the data, the investigator may oblige.

4.3.2 Metadata

The value of speech data can be determined according to the quality of metadata obtained. It is imperative to maintain metadata of the entire data collection for linguistic analysis.

After the recording is taken from the informant, personal details are collected. Care should be taken so that the signature and other formalities are completed as required.

The metadata of the speech corpus is made through the personal details taken from the informants. A typical copy of metadata sheet contains information as noted below:

Informant Data:

Name:
 Dataset ID:
 Address:
 Gender:
 Age Group: (with three options of 16 to 20, 21 to 50, and 50+)
 Educational Qualification: (with three options of School/Bachelors/Masters)
 Place of Elementary Education:
 Mother Tongue:
 Dialect (if any):

Investigator Data:

Name:
 Date:
 Place:
 Region:
 Environment:

It is to note that the name and the address of the informants are discarded while archiving metadata to keep the confidentiality and anonymity.

Dataset ID: It is a unique ID given to each speaker.

The following fields are considered for the distinctiveness of each data item recorded. Each field contributes certain features which pave way for diverse research.

Gender: The Speech data is taken from both male and female to capture the difference in intensity and pitch. The difference in vocal folds size between men and women makes them different in their pitched voices. Male voice usually has low pitch whereas a female voice is of high pitch. Pitch and intensity are proportional to each other.

Age Group: Different age groups exhibit difference in pitch and loudness. As the human body ages, it undergoes changes such as lessening strength, slower movements, degeneration of body tissues etc. these factors impact the voice as well. As people age their speech slows down, syllables and words are elongated, sentences are punctuated with more pauses for air. Scientific studies also show that as male and female age, the changing larynxes changes pitch and intensity. Age also affect the hearing process, which may make a person speak louder.

Educational Qualification: This determines the fluency and speed of reading speech data.

Place of Elementary Education: This parameter determines the effects of environment and dialect of a particular place of childhood which impacts the articulation of the speech.

Mother Tongue: Mother Tongue is one of the influential factors of a native speaker, for example In Karnataka, mainly in Canara region; it can be observed that the mother tongue of native Kannada speakers may be Tulu, Konkani, Chitpavani etc. This influences the articulation of Kannada speech in these areas.

Place: Place gives better information about the speech data collected. For example, Kannada spoken in Kundapura has its own distinct variety even when it belongs to Canara region.

Date: Date describes the timeline of data collected. It becomes useful information for historic research and language evolution in time line. It also dates the technology being used in that age.

Region: Region is an influencing factor of the language. Hence keeping the information about it with the data is always useful.

Environment: The recording environment information's like Indoor, Outdoor, School, Office, etc is useful, and its marking could be helpful in determining the noise level and the kind of noise that can be expected with the data.

Each of the datasets released contain a metadata sheet which has information about each of the audio files. A total of 25 fields are there in the metadata sheet.

A brief of each of these 25 fields/legends is given in the table below:

SL	Legend	Description
1	Language	Name of the Language
2	SpeakerID	Each speaker has a unique identity language. However, this is within the language. If one is working on speech corpus from more than one language, the IDs may get repeated.
3	ContentType	This corresponds to the notation of the content types noted above.
4	ContentID	This corresponds to the ID of the text being read out.
5	Gender	Notes gender, whether it is male, female or other.
6	AgeGroup	Three age groups of 16 to 20, 21 to 50, and 50+
7	Dialect	Notes the dialect of the language. An attempt has been made to cover all the dialects of the language as agreed upon in the academia of the language experts and linguists.
8	ReadInScript	The script in which the content has been provided to read in.
9	RecordingEnvironment	A brief info on the environment in which the recording has been done.
10	PowerSource	The source of the power using which the recording was done. It may be Li-ion, NiCd or Alkaline batteries.
11	RecorderManufacturer	Manufacturer of the recorder.
12	RecorderType	Type of the recorder. It is mostly 24-bit Linear PCM (R-09).
13	SamplingFrequency	Sampling frequency. It's mostly 48.
14	BitPerSample	Bit per sample. It is mostly 16-bit.
15	Channel	How many channels. All of LDC-IL data is stereo. Only data set is mono which is segregated and constitutes a separate dataset of its own.
16	State	Name of the Indian state/province to which the speaker belongs to.
17	District	Name of the Indian district to which the speaker belongs to.
18	Place	Name of the place to which the speaker belongs to.
19	MotherTongue	Mother tongue of the speaker. It is note that data has been taken from people who professo to speak the language. However, it may be that the speaker uses the target language as a second or third language. However, as long as the speaker confidently says (and it is also verified by other speakers of the community), some samples have been taken from these types of users as well. This adds to the variety of the speech data collected.
20	EducationalQualification	Highest educational qualification of the speaker.
21	PlaceOfElementaryEducation	Place of the elementary education. This usually corresponds to the early childhood experiences which happens to more than often affect the way a language spoken.
22	RecordingDate	Date when the recording took place.
23	Investigator	Name of the Investigator.
24	RecordedText	Text of the recorded speech (in the script of the language).
25	TextInRoman	Text of the recorded speech (in the Roman transliteration as per the LDC-IL transliteration scheme. If the text is long (as is the case with T1 and T2 content types), a reference of the corresponding file is given.)

Table 5: Metadata Legends and their Description

4.3.3 Data Transferring and Storing

After the data is collected for the day or when the SD card is full, the data needs to be transferred to the PC. It is important, to take certain precautions in this process so that the data is safely transferred. The data should be copied and pasted in the PC rather than cut and pasted. After successful transfer and rechecking the copied data, the SD card can be cleared.

For easier maintenance and organization of the data in PC, folder system is recommended for saving data. Each recorded wave file has to be labelled with corresponding speaker ID.

The investigator should try to get the required number of speakers/data before completing the field work within their schedule.

4.3.4 Observations

One of the reasons for error prone reading could be the over consciousness of the informant about being voice recorded. Despite being informed, the informant may try to read the data in a dramatic way, but may eventually lead to normal reading after few sentences. Even after the informants give consent and their data, they may later change their mind or express their concern about the text they have read. Some may even request to discard their recordings. In such cases, the investigator has to reassure them about their given data. If they still want their data to be discarded, they have to be accommodated. It is preferable to provide complete information to the informants to avoid such situations. In many instances informants assume that they are giving auditions for Radio Jockey vacancies, or some reality shows. They should be briefed about the purpose of data collection beforehand to avoid such situations.

The investigator may be in not so hospitable environments depending upon the region they are visiting. Proper precaution and aid is to be acquired before visiting such places.

The investigator may have to face challenges in food and accommodation since he/she travels in unfamiliar places. It is recommended to be prepared for such situations. The investigator should be prepared for all such hardships and take proper measures to minimize them beforehand.

5 ORGANISING AND ARCHIVING THE DATA

After the field work is completed, the data has to be stored in a server as soon as possible for safe keeping. Taking a backup of the saved data is also recommended as the data collected is of vital importance.

5.1.1 Text - Speech Mapping and Naming Conventions

After the data is stored, it is segmented and mapped with its corresponding text and metadata. Each recording 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 recorded data is shown bellow.

"LDC-IL_Scheduled_Kannada_Female_16To20_Contemporary Text-T1_SP-0035_T1-0035.wav"
"LDC-IL_Scheduled_Kannada_Male_21To50_Sentence-S_SP-0001_S-0004.wav"

Wave Surfer, free software, is used for segmentation of LDC-IL Speech data. It is an open source tool which can be downloaded freely from the web. While segmenting the speech data file for archiving, the introduction, content headings and any unnecessary speech are discarded. Only the dataset content is retained.

The ASR data is prepared keeping in the view, the stochastic systems such as HMMs or neural networks that do not use explicit rules for speech recognition. On the contrary, they rely on stochastic models which are trained using some statistical optimization procedure, with very large amounts of speech corpus.

5.1.2 Observations

While segmenting a single large recording containing all the content types, there may be instances where an informant has made an error and later corrected it. In such cases, it is always a good approach to segment a recording from the end of the file in reverse order so that the correct utterance can be found before incorrect utterance; hence the incorrect utterance can be discarded/ignored. One may observe the interventions of investigator or other people between reading two data items which may also need to be discarded.

6 DATA VERIFICATION AND QUALITY CONTROL

Since mapping audio recordings with its corresponding text and other metadata information is a manual task. The process is prone to human errors; the data verification process will be done

Much of the audio text mapping is automated, but in case of distinct set texts, and other metadata entry is done by human needs verification. In this,

- The Audio recording of each speaker is checked against the mapped text.
- Each distinct text audio recording will be matched with automated entries of the same speaker to check for any mis-mapping of speaker.
- Metadata like Gender, age group, District etc are selective part of manual data entry and could be prone to errors so verification is needed.
- Metadata like Dialect entry, place, etc are keyed in by manual data entry and could be prone to errors like typo errors so verification needs to be done.
- The audio segments could be duplicated because of system/network errors and need to be checked.
- At the time of data segmentation, one might have saved the whole file instead of selected part. Such cases need to be checked.
- Some audio segments may not get migrated to the system because of wrong naming conventions. Such segments will be handpicked and corrected and migrated into the system.

7 MANIPURI RAW SPEECH CORPUS

7.1 INTRODUCTION

India, one of the most linguistically diverse countries, has five language families, namely – the Indo-Aryan, the Dravidian, the Austro-Asiatic, the Tibeto-Burman and the Andamanes respectively. It is in good health that the language policy of India is elucidated in the Constitution, implemented through various executive orders issued from time to time as well as the judicial pronouncement since 1950 focusing on the scope of being language-development oriented and language-survival oriented. In this parlance was the inclusion of the Eighth Schedule in the constitution providing formal and constitutional recognition to dominant regional languages in the sphere of administration, education, economy and social status.

7.1.1 Manipuri in the Eighth Schedule of the Constitution of India

Manipuri languages obtained its due place and recognition in the Constitution of India being included in the Eight Schedule, according to the Seventy First Amendment of the Constitution, on the 20th August, 1992. Here it is noteworthy to mention that Manipuri language, the state official language in Manipur, is also the first Tibeto-Burman (henceforth TB) language included in the said Eighth Schedule of the Constitution of India. Another mentionable point here is that it is the only language amongst the TB-languages in India which has its own scripts and written literature.

Linguistically speaking, Manipuri is the lingua-franca, and also an Inter-Tribal language amongst the speakers of different dialects and other minor languages inhabiting both the Hill and Valley areas of the state of Manipur.

Adopted as the medium of instruction and examination from Primary School level to College and University level, Manipuri has been offered as a subject of study and research not only by the Meetei native speakers alone but also by other community groups in the state of Manipur.

7.1.2 Area and Population of Manipur

Manipur has area of 22,327 sq. kms. According to census of 1981, the biggest valley area of the state, now known as the Imphal Valley is about 1,843 sq. kms which is roughly 9% of the total area of the entire state. Physically, Manipur comprises of two parts, the hills and the valley. The valley lies at the centre surrounded by hills on all sides. The hills cover about 9/10 of the total area of the State. The State has a population of almost 3 million, including the Meeteis, who are the majority ethnic group in the state and other ethnic community groups who speak a variety of Tibeto-Burman languages. In fact, the term Meetei is an endonym or autonym while Manipuri is an exonym. Manipur lies between latitude 23050' and 25030' North, and longitude 93010' and 94030' East. Having an oval shape area on the basis of geographical position of the state on the surface of the Earth, longer in north and south, and shorter in east and west, in length, the state enjoys mild sub-tropical temperate monsoon climate with temperature varying from a little above 00C and below 350C.

7.1.3 A cursory Glance at Manipuri Dialects

It is a generally held view that there was no particular language variety called Meeteilon (now known as Manipuri in the constitution) in its pre-historic period. What is believed so far is that there were different varieties in the form of dialects of the same speech spoken and used in different parts of the state.

Nevertheless, there is a mutual intelligibility between these dialects in which speakers of different but related dialects can readily understand each other without prior familiarity or any special effort. Gradually, a particular variety came into existence taking its certain uniform form along with typical linguistic features. This particular variety is the standard Meeteilon/Manipuri, which is the by-product of different linguistic forces and different dialectal elements contributed to it. Henceforth, the dialect of the Kangla Imphal became its base while other dialects of the same speech continued to exhibit shares in it. Today, with cycle of time, we can observe that the intelligibility lying between Imphal dialect and other dialects such as Kakching, can be asymmetric in such a way that speakers of one dialect, say Kakching dialect, happen to understand more of the other, the Imphal dialect, than the speakers of Imphal understanding Kakching dialect.

Manipuri (locally known as Meeteilon) has dialects of Kakching, 45 kms far away from Imphal, located in the south eastern part of the state, Awang Sekmai, 17 kms from Imphal, located western part of the state, Andro, 25 kms from Imphal, located in the eastern part of the state, Phayeng, 16 kms from Imphal, located in the western part of the state, Kwatha, 102 kms from Imphal, located in the eastern most part of the state, Thangga, 48 kms from Imphal, located in the western part of the state, etc. In this way, there is a dialect continuum or dialect chain for this Manipuri language typically occurring in a long-settled Meetei population. Standard Imphal variety within such a dialect continuum was developed and codified serving as an authority for part of it across the various geographical areas of the monolingual community. Imphal dialect is now used for official purposes, heard on radio and televisions and considered the standard form of their speech so that any standardizing changes in their speech are always oriented towards that Imphal variety. In these apparent cases, other local dialectal varieties are said to be dependent on, or heteronomous with respect to, the standard Imphal variety. This is how the formation of dialects took place and Manipuri has a standard Imphal variety, together with its dependent varieties called “dialects” of the language even though this standard variety is mutually intelligible with the other rest of the dialectal forms from the continuum.

7.1.4 Intra-regional speech variation and LDC-IL dataset

Language variation exists even in monolingual communities. The sociolinguistic elements such as social status, gender, age and ethnicity etc. happen to get reflected in the language people speak and happen to turn out to be important dimensions of the speakers’ identity in their community. Every dialect has its own unique linguistic features which the group shares with each other within the small group. It is a fact that no two people speak exactly the same exhibiting infinite source of variation in their speeches. One can notice that a sound spectrograph shows that even a single vowel could be pronounced in hundreds of minutely different manners, most of which the listeners cannot even register. However, there are certain common features of speech which each dialect exhibits and this feature is shared by the group concerned differentiating them from the other group which again has its own common one. In the present

scenario of Manipuri dialectal variations, the pronunciation, grammar, and vocabulary of Kakching or Awang Sekmai speakers of Manipuri are in some respects found quite distinct from that of people from Imphal. Since the standard Imphal dialect is always the first to be codified, the act of defining other dialects is done through contrasting them with the standard. In this way, one can capture how Kakching or Awang Sekmai dialect features contrast with the standard Imphal dialect features. In this perance, certain linguistic features identifying regional tones and intonations, phonemic distributions (as observed in the dataset of Phonetically Balanced vocabulary-W4 and Phonetically Balanced Sentences-S, various pronunciations reflected in both regional and non-regional vocabulary items such as person names and place names (as observed in the dataset of Person names W2 and Place Names W2) etc., have been well housed based on a standard parameter of LDC-IL dataset. The LDC-IL Manipuri Raw Speech Corpus reflects the speech varieties of the same language providing the speakers' social background. In a nutshell, such speech corpus can be empirically used for developing NLP tools such as speech synthesis, speech recognition, spoken language systems, and speaker recognition /verification, etc. which generally employs algorithms working with acoustic and language modelling. On the other hand, this speech corpus can also be used for research purposes such as phonetic research, sociolinguistic research, psycholinguistics research and other language acquisition investigation purposes. Since the same dataset is applied to all the said regional dialects, vocabulary and grammar cannot be distinguished from each other; however, there are clear-cut distinctions in pronunciations observed amongst the respective dialects.

In its real sense, such language-specific resources like Manipuri Raw Speech Corpus can be used for building Manipuri TTS systems since it can be viewed as providing (i) audio recordings (ii) pronunciation lexicon and (iii) phone sets containing phonetic features for each phoneme. Since the corpus is in its raw status, the researchers/experts can specify the transcripts to the recorded segments; give letter-to-sound rules and develop phone sets that contain phonetic features for each phoneme, which will finally be used as the specific features in models of the spectrum and prosody. In fact, this Manipuri Speech Corpus is the actual linguistic resources and data for such speech synthesis task.

As one can notice that the LDC-IL Manipuri Raw Speech corpus has a total number of 620 speakers (310 Female and 310 Male) collected from three regional dialects, namely Imphal, Kakching, and Awang Sekmai respectively, the whole database can be said to be consisting of speeches with various voice characteristics and speaking styles which are more feasible features for the speech synthesis to achieve maintaining some acceptable degree of naturalness.

8 DATASET PREPARATION FOR MANIPURI

For the selected dialects such as Imphal, Kakching and Awang Sekmai, LDC-IL prepared the following dataset for which the prompt sheets were prepared.

DATASET CONTENTS FOR MANIPURI

Content Type	Count
Creative Text	16
Date	2
Command and Control Words	249
Most Frequent Words	1000
Person Name	501
Place Name	324
Sentences	200

Table 6: LDC-IL Speech Dataset

Distinct News Items are prepared to get the audio recording of contemporary text along with the aforementioned content types.

Each prompt sheet that consists of a distinct news item and selected part of the dataset is prepared as follows:

Content Type	Content that Each typical prompt sheet had	Content selection type
Contemporary Text	1 Text	Distinct Text
Creative Text	1 text	Random Text selected from dataset*
Sentences	25 Sentences	Random set selected from dataset*
Command and Control Words	30 Words	Random set selected from dataset*
Person Names	20 Words	Random set selected from dataset*
Place Names	10 Words	Random set selected from dataset*
Most Frequent Words	30 Words	Random set selected from dataset*
*randomly selected by machine		

Table 7: Table of Contents in LDC-IL Dataset

Once all these preparations are made, the investigators start collecting the data. The collection of data is carried out in four phases:

Region	Year	Field Investigator
Imphal	2008	Amom Nandaraj Meetei
Kakching	2009	Amom Nandaraj Meetei
Awang Sekmai	2010	Amom Nandaraj Meetei
Imphal	2013	Yumnam Premila Chanu

Table 8: Four Phases of Speech Data Collection

9 TRANSLITERATIONS IN LDC-IL MANIPURI READ CORPUS

For easy reference and uniformity, the recorded text in the metadata file, is transliterated from Bengali to Roman letters. Numeric characters were transliterated from Bengali to Hindu-Arabic system.

The LDC-IL transliteration scheme of Manipuri (in Bengali scripts) to Roman is given below.

LDC-IL Transliteration Schema													
Manipuri characters to Roman and Manipuri Numerals to Hindu-Arabic													
Scripts	Vowels and Vowel Signs in Bengali Scripts												
Bengali	অ	আ	ই	ঈ	উ	ঊ	ঋ	এ	ঐ	ও	ঔ		
Bengali	া		ি	ী	ু	ূ	্	ে	ৈ	ো	ৌ	ং	ঃ
Roman	a	A	i	I	u	U	x	E	ai	O	au	M	H
Consonants							Unreleased Consonants						
Bengali	ক	খ	গ	ঘ	ঙ	ক	ক্		ঙ্				
Roman	ka	kha	ga	gha	ng'a	ka	k		ng'				
Bengali	চ	ছ	জ	ঝ	ঞ	চ	ম্		প্				
Roman	ca	cha	ja	jha	nj'a	ca	m		p				
Bengali	ট	ঠ	ড	ঢ	ণ	ট	ন্	ত্	ল্				
Roman	Ta	Tha	Da	Dha	Na	Ta	n	t	l				
Bengali	ত	থ	দ	ধ	ন	ত							
Roman	ta	tha	da	dha	na	ta							
Bengali	প	ফ	ব	ভ	ম	প							
Roman	pa	pha	ba	bha	ma	pa							
Bengali	য	র	ল	শ	স	ষ	হ	ড়	ঢ়	য়	ৎ		
Roman	ya	ra	la	sha	Sa	sa	ha	D'a	Dh'a	Ya	t		
Numerals (Bengali to Hindu-Arabic)													
Bengali	০	১	২	৩	৪	৫	৬	৭	৮	৯			
Hindu-Arabic	0	1	2	3	4	5	6	7	8	9			

The following citation visualizes the running idea.

ContentID: S-0001

Recorded Text: মহাক্ষা পুনগি অসেঙবা ইসিঙ তপ্পা থক্লি।

Transliteration: mahAkna punagi aseng'abA ising'a tapna thakli.

10 SUMMARY OF THE CORPUS

In the sections below, we provide the tabular details of the different content types of the Manipuri Raw Speech Corpus on the basis of various yardsticks being filtered out from the metadata sheet. These figures may be helpful in tuning the corpus for various purposes of training, testing and evaluating various algorithms as well as providing useful insights into the dataset. The data size is of total duration 156:28:32 (hh:mm:ss) comprising 66,231 audio segments.

10.1 SUMMARY OF THE AUDIO SEGMENTS

The total number of Audio Segments along with their distribution in terms of Gender and Age for the Manipuri Speech Dataset is shown below.

LDC-IL Manipuri Speech Data Status	Gender →	Female			Male		
	Age Group →	16-20 Years	21-50 Years	50+ Years	16-20 Years	21-50 Years	50+ Years
Content Type	Total Segments	Segments	Segments	Segments	Segments	Segments	Segments
Contemporary Text (News)-T1	530	29	187	55	26	170	63
Creative Text-T2	588	29	194	72	34	187	72
Sentence-S	10979	600	3277	1600	600	3301	1601
Date-D	866	48	257	128	46	261	126
Command and Control Words- W1	13129	714	3928	1919	720	3928	1920
Person Name-W2	8789	481	2625	1280	480	2641	1282
Place Name-W2	4394	240	1311	640	241	1321	641
Most Frequent Word-Part-W3A	13167	722	3929	1920	720	3956	1920
Most Frequent Word-FullSet- W3B	6992	1000	1996	0	998	1998	1000
Phonetically Balanced-W4	4518	753	753	753	753	753	753
Form and Function Word- W5	2279	380	380	380	380	379	380

Table 9 : Manipuri Audio Segments and their Distribution

10.2 DURATION OF THE RAW SPEECH DATA

The table below shows the duration of each of the content type and their distribution across the gender and age factors.

LDC-IL Manipuri Speech Data Status	Gender →	Female			Male		
	Age Group →	16-20 Years	21-50 Years	50+ Years	16-20 Years	21-50 Years	50+ Years
Content Type	Total Duration	Duration (hh:mm:ss)	Duration (hh:mm:ss)	Duration (hh:mm:ss)	Duration (hh:mm:ss)	Duration (hh:mm:ss)	Duration (hh:mm:ss)
Contemporary Text (News)-T1	59:47:22	02:55:21	22:46:16	04:55:44	03:31:13	20:26:26	05:12:22
Creative Text-T2	53:59:35	02:15:42	20:12:16	04:13:38	03:29:04	19:14:49	04:34:06
Sentence-S	10:01:41	00:33:51	03:03:27	01:25:40	00:31:00	02:56:09	01:31:34
Date-D	1:12:04	00:04:20	00:21:03	00:10:27	00:04:00	00:21:55	00:10:19
Command and Control Words-W1	8:00:02	00:26:41	02:26:28	01:10:38	00:24:11	02:19:52	01:12:12
Person Name-W2	7:14:04	00:24:22	02:13:58	01:03:12	00:21:39	02:05:45	01:05:08
Place Name-W2	2:46:29	00:09:34	00:51:09	00:24:34	00:08:05	00:48:13	00:24:54
Most Frequent Word-Part-W3A	6:48:50	00:22:50	02:03:59	00:59:48	00:20:23	01:59:36	01:02:14
Most Frequent Word-FullSet-W3B	2:48:42	00:26:26	00:43:29	00:00:00	00:21:08	00:50:36	00:27:03
Phonetically Balanced-W4	2:25:53	00:27:41	00:19:54	00:28:09	00:26:27	00:19:45	00:23:57
Form and Function Word-W5	1:23:50	00:14:10	00:12:07	00:12:09	00:14:18	00:19:07	00:11:59

Table 10: Duration of the Collected Manipuri Speech Data

10.3 DISTINCT SET

The Distinct Set usually contains data which is distinct to each speaker and is rarely repeated. The LDC-IL speech data set contains newspaper extracts which are read by each speaker.

10.3.1 The Contemporary Text (News) - T1

Distinct Text Extracts from Newspapers are recorded from the informants to get the speech data of Contemporary Text. The distribution of data is as follows:

Age Group	Total Audio Segments	Gender-wise Distribution		Region-wise Distribution					
				IMPHAL		KAKCHING		SEKMAI	
		Female	Male	Female	Male	Female	Male	Female	Male
16 To 20	55	29	26	13	10	9	8	7	8
21 To 50	357	186	171	100	84	46	42	41	44
50+	118	55	63	13	19	21	24	21	20
Total	530	270	260	125	114	76	74	69	72

Table 11: Distribution of Manipuri Contemporary Text (News) Data

10.4 RANDOM SET

The Random Set data contains content types sampled by machine for each speaker. They are basically sampled from the collection of master data sets available. The random sets are given below:

10.4.1 The Created Text-T2

One randomly selected text of literature out of 16 texts from the prepared dataset is recorded from the informants to get the speech data of Created Text. The distribution of data is as follows:

Age Group	Total Audio Segments	Gender-wise Distribution		Region-wise Distribution					
				IMPHAL		KAKCHING		SEKMAI	
		Female	Male	Female	Male	Female	Male	Female	Male
16 To 20	63	29	34	13	18	8	8	8	8
21 To 50	381	193	188	107	101	46	42	41	44
50+	144	72	72	30	28	21	24	21	20
Total	588	294	294	149	148	75	74	70	72

Table 12: Distribution of Manipuri Created Text: T2

10.4.2 The Sentences-S

The Sentences-content type consists of a list of sentences that can be representative of which all the phonemes in Manipuri can occur in various positions. 25 Randomly selected Sentences are recorded from a list of 208 sentences. The distribution of data is as follows:

Age Group	Total Audio Segments	Gender-wise Distribution		Region-wise Distribution					
				IMPHAL		KAKCHING		SEKMAI	
		Female	Male	Female	Male	Female	Male	Female	Male
16 To 20	1200	600	600	200	200	200	200	200	200
21 To 50	6578	3277	3301	1100	1150	1151	1051	1026	1100
50+	3201	1600	1601	550	500	525	601	525	500
Total	10979	5477	5502	1850	1850	1876	1852	1751	1800

Table 13: Distribution of Manipuri Sentences

10.4.3 The Date-D

The answer to one randomly selected question from the list of 2 questions is recorded to get the Date Format of the informants. The distribution of data is as follows:

Age Group	Total Audio Segments	Gender-wise Distribution		Region-wise Distribution					
				IMPHAL		KAKCHING		SEKMAI	
		Female	Male	Female	Male	Female	Male	Female	Male
16 To 20	94	48	46	16	16	16	14	16	16
21 To 50	518	257	261	88	92	87	82	82	87
50+	254	128	126	44	40	42	46	42	40
Total	866	433	433	148	148	145	142	140	143

Table 14: Distribution of Manipuri Date Format

10.4.4 Command and Control Words-W1

The command and control words content type consists of a list of 187 words that is a representative of which most of the command and control words occur in Manipuri. 30 randomly selected words are recorded from a list of words. The distribution of data is as follows:

Age Group	Total Audio Segments	Gender-wise Distribution		Region-wise Distribution					
				IMPHAL		KAKCHING		SEKMAI	
		Female	Male	Female	Male	Female	Male	Female	Male
16 To 20	1434	714	720	235	240	240	240	239	240
21 To 50	7856	3928	3928	1320	1380	1379	1259	1229	1289
50+	3839	1919	1920	659	600	630	720	630	600
Total	13129	6561	6568	2214	2220	2249	2219	2098	2129

Table 15: Distribution of Manipuri Command and Control Words

10.4.5 Person Names –W2

The Person Names content type consists of a list of 500 popular Pan Indian and regional person names. 20 randomly selected names are recorded from a list of names. The distribution of data is as follows:

Age Group	Total Audio Segments	Gender-wise Distribution		Region-wise Distribution					
				IMPHAL		KAKCHING		SEKMAI	
		Female	Male	Female	Male	Female	Male	Female	Male
16 To 20	961	481	480	162	160	159	160	160	160
21 To 50	5266	2625	2641	880	920	920	839	825	882
50+	2562	1280	1282	440	400	420	481	420	401
Total	8789	4386	4403	1482	1480	1499	1480	1405	1443

Table 16: Distribution of Manipuri Personal Names

10.4.6 Place Names-W2

The Place Names content type consists of a list of 324 popular Pan Indian and regional place names. 10 randomly selected names are recorded from a list of names. The distribution of data is as follows:

Age Group	Total Audio Segments	Gender-wise Distribution		Region-wise Distribution					
				IMPHAL		KAKCHING		SEKMAI	
		Female	Male	Female	Male	Female	Male	Female	Male
16 To 20	481	240	241	80	80	80	80	80	81
21 To 50	2632	1311	1321	440	460	460	420	411	441
50+	1281	640	641	220	200	210	241	210	200
Total	4394	2191	2203	740	740	750	741	701	722

Table 17: Distribution of Manipuri Place Names

10.4.7 Most Frequent Words-PART-W3A

The Most Frequent Words-Part content type consists of a list of 1000 most frequent words. 30 randomly selected words are recorded from a list of such frequent words. The distribution of data is as follows:

Age Group	Total Audio Segments	Gender-wise Distribution		Region-wise Distribution					
				IMPHAL		KAKCHING		SEKMAI	
		Female	Male	Female	Male	Female	Male	Female	Male
16 To 20	1442	722	720	242	240	240	240	240	240
21 To 50	7885	3929	3956	1320	1379	1380	1254	1229	1323
50+	3840	1920	1920	660	598	630	722	630	600
Total	13167	6571	6596	2222	2217	2250	2216	2099	2163

Table 18: Distribution of Manipuri Most Frequent Words – Part

10.5 FULL SET

The Full Sets are the master set of certain data sets which are read completely by few selected speakers in each group. The full sets are as below:

10.5.1 Most Frequent Words-Full-W3B

The Most Frequent Words-Full content type consists of a list of 1000 most frequent words. In this case, all the 1000 words are recorded from the informants. The distribution of data is as follows:

Age Group	Total Audio Segments	Gender-wise Distribution		Region-wise Distribution			
				IMPHAL		KAKCHING	
		Female	Male	Female	Male	Female	Male
16 To 20	1998	1000	998	0	0	1000	998
21 To 50	3994	1996	1998	1000	1000	996	998
50+	1000	0	1000	0	0	0	1000
Total	6992	2996	3996	1000	1000	1996	2996

Table 19: Distribution of Manipuri Most Frequent Word-Full

10.5.2 The Phonetically Balanced Words-W4

The Phonetically Balanced Words type contains a list of words where most of the phones of Manipuri language have occurred in all the possible positions of a word. In this case, all the 390 words are recorded from the informants in such a way that they utter those words three times. The distribution of data is as follows:

Age Group	Total Audio Segments	Gender-wise Distribution		Region-wise Distribution			
				IMPHAL		KAKCHING	
		Female	Male	Female	Male	Female	Male
16 To 20	1506	753	753	374	374	379	379
21 To 50	1506	753	753	374	374	379	379
50+	1506	753	753	374	374	379	379
Total	4518	2259	2259	1122	1122	1137	1137

Table 20: Distribution of Manipuri Phonetically Balanced Words-W4

10.5.3 The Form and Function Words-W5

The Form and Function Words content type contains a list of 432 words which is a representative of which most of the form and function words occur in Manipuri. All the words are recorded from the informants in such a way that they utter those words three times. The distribution of data is as follows:

Age Group	Total Audio Segments	Gender-wise Distribution		Region-wise Distribution			
				IMPHAL		KAKCHING	
		Female	Male	Female	Male	Female	Male
16 To 20	760	380	380	189	189	191	191
21 To 50	759	380	379	189	189	191	190
50+	760	380	380	189	189	191	191
Total	2279	1140	1139	567	567	573	572

Table 21: Distribution of Manipuri Form and Function Words-W5

10.6 NATIVE SPEAKERS DISTRIBUTIONS

The following table displays the overall distributions of the native speakers who read the data during the fieldworks undertaken at different three regions of the state of Manipur.

Age Group	Total Speakers	Gender-wise Distribution		Region-wise Distribution					
				IMPHAL		KAKCHING		SEKMAI	
		Female	Male	Female	Male	Female	Male	Female	Male
16 To 20	74	35	39	15	20	12	11	8	8
21 To 50	393	199	194	109	104	49	45	41	44
50+	154	76	78	32	31	23	27	21	20
Total	620	310	310	156	155	84	83	70	72

Table 22: Distribution of Manipuri Native Speakers

10.7 CONCLUSION

This documentation is representative of contemporary speech corpus for Manipuri in response to the corpus generation revolution undertaken by the government of India for the development of Indian scheduled languages in technological media world which was initiated in the form of the technological development works on scheduled languages in 1991. LDC-IL has created Manipuri Raw Speech Corpus, which is the backbone of automatic speech recognition and also for different speech recognition tasks in Manipuri language. The database comprises of appropriate words, sentences and paragraphs spoken by the typical users in realistic acoustic or natural environments following the data collection guidelines of LDC-IL.

10.8 REFERENCE

- Amom, N.M., Sakkan Th, Sarma, A. 2013. CORPUS Building of Endangered Languages, in *Language Endangerment in South Asia*, (Vol-I), (eds.) by Ganeshan, M, et al, Centre of Advanced Study in Linguistics, Annamalai University, Annamalai nagar.
- Dash, N.S. 2005. *Corpus Linguistics and Language Technology*. New Delhi: Mittal Publications.
- Dash, N.S. 2006. *Speech corpora Vs. Text Corpora: Need for Separate Development*. Indian Linguistics. Vol. 67. Nos. 1-4. Pp. 65-82.
- Leech, G. 1991. "The state of the art in corpus linguistics". In, Aijmer, K. and B. Altenberg (Eds.) *English Corpus Linguistics: Studies in Honour of Jan Svartvik*. London: Longman. Pp. 8-29.