

Zeliang Parallel Text Corpus: Linguistic Features and Structures



ZELIANG PARALLEL TEXT CORPUS: LINGUISTIC FEATURES AND STRUCTURES



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.

Authors:

***Dr. Kamaraj S,
Dr. Rejitha K. S,
Dr. Narayan Choudhary,
Prof. Shailendra Mohan***

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

भारतीय भाषा संस्थान

उच्चतर शिक्षा विभाग, शिक्षा मंत्रालय, भारत सरकार, मानसगंगोत्री, हुन्सूर रोड, मैसूरु-570006, कर्नाटक

CENTRAL INSTITUTE OF INDIAN LANGUAGES

Department of Higher Education, Ministry of Education, Government of India
Manasagangotri, Hunsur Road, Mysuru-570006, Karnataka

www.ciil.gov.in :: e-✉ ciil.publication@gmail.com :: ☎ +91-821-234-5182

Zeliang Parallel Text Corpus: Linguistic Features and Structures

Authors: Dr. Kamaraj S, Dr. Rejitha K. S, Dr. Narayan Choudhary, Prof. Shailendra Mohan

Contributors: Lungchiepyile, Machipeung Thou, Kangzangyile Khate

e-ISBN:

CIIL Publication No.:1710

First published: AD 2026 March: Phalguna 1947 Śaka

© Central Institute of Indian Languages, Mysuru 2026

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: Dr.Kamaraj S

Table of Contents

Zeliang Parallel Text Corpus	1
1.1 LDC-IL Parallel Text Corpus.....	1
1.2 Challenges in Creating A Parallel Corpus.....	2
1.3 Zeliang: A Brief Introduction.....	2
1.4 Zeliang parallel Corpus	3
1.5 Summary of Zeliang parallel Corpus	4
1.6 Reference.....	4

ZELIANG PARALLEL TEXT CORPUS

Dr.Kamaraj S

1.1 LDC-IL PARALLEL TEXT CORPUS

LDC-IL has developed parallel corpus in Indian mother tongues. India has 270 mother tongues as per 2011 census. Most of these mother tongues have no written text readily available and they are mostly oral languages. Most of them do not have any script of their own or use one or more than one dominant script used in that region. Following the requirements of the NEP-2020, LDC-IL initiated to collect sentences from different mother tongues that could give an idea about the basic structure of these 270 mother tongues.

A parallel corpus is a collection of texts in two or more languages that are carefully aligned sentence wise in pairs and placed either side by side or matched with a common sentence/phrase identification number. It is widely used in linguistics, translation studies, and Natural Language Processing. Researchers use it to compare vocabulary, grammar, and meaning across languages, while computational systems rely on it to train statistical and neural machine translation models. Parallel corpora also support language learning by providing real translation examples and enable cross-language information retrieval, allowing users to access information in multiple languages efficiently.

LDC-IL developed a framework to accommodate all the linguistic features of mother tongues. The sentences were collected from Indian scheduled languages especially from descriptive grammar books. The sentences cover the general grammatical and structural types like adjectival phrases, adverbial phrases, compounding, Interrogative sentences, Imperative sentences, different speech types, statements, structural question, coordination, subordination and so on. Moreover, it includes linguistic features such as use of adposition, anaphora, derivational morphology, inflectional morphology and language specific features like emphasis, equative, possession, reciprocation, reflexive expression and so on. For more details on the source collection for this parallel corpus, please refer to the [LDC-IL Corpus Insights \(2025\)](#).

This parallel corpus can be used in AI development by enabling the creation of Indic multilingual foundation models, develop cross-lingual knowledge graphs and create low-resource language adapters. By providing high-quality linguistic data, it strengthens natural language understanding and generation. As India rapidly advances in the digital age, it is essential to preserve and promote undocumented mother tongues by bringing them onto digital platforms, ensuring they are sustained and able to thrive for future generations.

Low-resource languages have a significant impact on both AI development and digital world. When AI systems are trained mostly on high-resource languages, they become biased toward those linguistic and cultural contexts. Including low-resource languages helps build stronger multilingual foundational models. Exposure to diverse grammar structures, vocabulary, and linguistic patterns improves cross-lingual learning and generalization. Language carries cultural knowledge. Incorporating low-resource languages enriches AI systems with diverse worldviews, idioms, and traditional knowledge systems. If a language is not supported digitally, its speakers

may be excluded from online education, e-governance, healthcare services, and digital commerce. When digital tools support local languages, more people can participate in the economy. Local languages often contain indigenous knowledge about agriculture, medicine, environment, and culture and therefore bringing them into digital ecosystem safeguards the linguistic treasures that language might contain.

Supporting low-resource languages makes AI more inclusive, intelligent, and representative of linguistic diversity. In India, where hundreds of mother tongues co-exist, investing in low-resource language technologies is essential for equitable digital growth.

1.2 CHALLENGES IN CREATING A PARALLEL CORPUS

Creating a Parallel Corpus for mother tongues in India involves several linguistic, technical, and socio-cultural challenges. The major difficulty is the lack of standardized orthography in many regional and tribal mother tongues. Most of the mother tongues use the prominent script used in their region. Some mother tongues may use multiple scripts or even a newly created writing systems, which complicates the corpus development. Another challenge is the scarcity of written and digitized resources. Many mother tongues are used for daily communication only.

Dialectal variation is another serious issue, as a single mother tongue may have numerous regional varieties with differences in vocabulary, pronunciation, and grammar. This makes it difficult to ensure consistency and representativeness in the corpus. Accurate translation is a concerned matter. Proficient bilingual speakers may be limited for low-resource languages. Maintaining semantic equivalence across mother tongues with different cultural contexts can be challenging. Sometimes equivalent words are not available due to the socio-cultural differences. These factors affect large-scale parallel corpus development in Indian mother tongues. The creation of reliable and balanced parallel corpora for Indian mother tongues is a complex and resource-intensive task.

1.3 ZELIANG: A BRIEF INTRODUCTION

Zeliang language (often referred to as Zeliang, representing the Zeme, Liangmai, and Rongmei varieties) is a Tibeto-Burman language cluster predominantly spoken in the states of Nagaland, Manipur, and Assam. It forms an important part of the cultural and ethnic identity of the Zeliangrong Naga community. According to the Census of India 2011, the combined speaker population of these varieties is over two lakh, though the numbers vary for each subgroup. Zeliangrong languages developed as distinct yet closely related speech forms within the Naga subgroup of the Sino-Tibetan language family. Historically, they functioned primarily as oral languages, preserving history, customary laws, myths, folktales, and ritual traditions through oral transmission. With the arrival of Christian missionaries in the 19th and early 20th centuries, these languages began to be written systematically using the Roman script. Today, the Roman script is widely used for writing Zeme, Liangmai, and Rongmei, with adaptations to represent their phonological features.

Zeliang literature initially emerged through translated religious texts such as the Bible, hymn books, and catechism materials. Over time, original writings including poetry, short stories,

folklore collections, and educational materials have been produced in the respective varieties. Although the written literary tradition is relatively recent compared to classical Indian languages, the oral literary heritage of the Zeliangrong community is rich and culturally significant. The language cluster has deep historical roots within the Naga cultural framework and continues to hold contemporary relevance among its speakers. It is actively used in daily communication, church services, community events, and cultural festivals. Local publications, community organizations, and educational initiatives contribute to its continued vitality. With its unique linguistic structure, oral tradition, and strong cultural identity, the Zeliang language cluster occupies an important place in the linguistic diversity of Northeast India. Ongoing efforts in documentation, literacy development, and digital preservation are helping to sustain and promote the language for future generations.

1.4 ZELIANG PARALLEL CORPUS

Zeliang language (commonly referred to as Zeliang, representing Zeme, Liangmai, and Rongmei varieties) is an agglutinative Tibeto-Burman language cluster with notable morphological and phonological features. Grammatical relationships are largely expressed through suffixation, including markers for tense, aspect, mood, number, and case-like functions. The language exhibits morphophonemic alternations during affixation, particularly in verb forms and nominal constructions. Word formation processes such as compounding and derivation are productive, and tone plays a meaningful role in distinguishing lexical or grammatical contrasts in certain varieties. Zeliang languages are written in the Roman script, introduced through missionary influence and later adapted to represent their phonemic systems. The orthography uses combinations of Roman letters to capture consonant clusters, vowel distinctions, and in some cases tonal features. Although efforts have been made toward standardization, minor orthographic variations may still exist across Zeme, Liangmai, and Rongmei varieties. The basic word order of Zeliang languages generally follows the Subject–Object–Verb (SOV) pattern. Postpositions are commonly used, and modifiers typically precede the nouns they qualify. Despite having a relatively smaller speaker base compared to major Indian languages, Zeliang is linguistically rich and structurally significant within the Tibeto-Burman family. Its tonal system, agglutinative morphology, and syntactic structures make it valuable for typological and comparative linguistic research.

Zeliang is also an important resource for developing parallel text corpora with other Indian mother tongues. A parallel corpus involving Zeliang can provide valuable insights into syntactic divergence between Tibeto-Burman, Indo-Aryan, and Dravidian languages, as well as into morphological complexity, code-mixing practices, and script variation. Such corpora are especially significant in multilingual contexts like India, contributing to research in language documentation, computational linguistics, machine translation, and the preservation of indigenous linguistic heritage.

This Zeliang parallel corpus constitutes a systematically translated dataset derived from an original English corpus, ensuring structured cross-linguistic alignment and semantic correspondence

1.5 SUMMARY OF ZELIANG PARALLEL CORPUS

The Zeliang parallel text corpus is aligned with English and 146 mother tongues of India. These mother tongues are Anal, Angami, Apatani, Are, Assamese, Awadhi, Bagheli/Baghel Khandi, Bagri Rajasthani, Bagri, Balti, Bengali, Bhadrawahi, Bharmauri/Gaddi, Bhojpuri, Bilaspuri Kahluri, Bodo, Brajbhasha, Bundeli/Bundel khandi, Chakru/Chokri, Chambeali/Chamrali, Chang, Chhattisgarhi, Chirri, Chungli, Churahi, Coorgi/Kodagu, Deori, Dhundhari, Dimasa, Dogri, Gangte, Garhwali, Garo, Gojri/Gujjari/Gujar, Gujarati, Gujar, Halabi, Handuri, Hara/Harauti, Haryanvi, Hindi Multani, Hindi, Irula/Irular Mozhi, Kabui, Kachchhi, Kangri, Kannada, Karbi/Mikir, Kashmiri, Khandeshi, Khari Boli, Khasi, Khezha, Khiemnungan, Khortha/Khotta, Kisan, Kodava, Kokbarak, Kolami, Koli, Kom, Konda, Konkani, Konyak, Koya, Kudubi/Kudumbi, Kuki, Kurmali Thar, Ladakhi, Lepcha, Liangmei, Limbu, Lotha, Lyngngam, Magadhi/Magahi, Maithili, Malayalam, Malvi, Manipuri, Mao, Mara, Maram, Marathi, Maring, Mech/Mechhia, Mewari, Mewati, Miri/Mising, Mishmi, Mizo, Mongsen, Monpa, Mundari, Muwasi, Nawait, Nepali, Nimadi, Nissi, Nocte, Odia, Pahari, Paite, Palmuha, Pania, Pawari/Powari, Phom, Pnar/Synteng, Pochury, Poula, Punjabi, Purkhi, Rai, Rajasthani, Reang, Rengma, Rongmei, Sadan/Sadri, Sambalpuri, Sangtam, Sanskrit, Santali, Saurashtra/Saurashtri, Sema, Shina, Sindhi, Sirmauri, Sugali, Surjapuri, Talgalo, Tamil, Tangkhul, Telugu, Thado, Tibetan, Tikhir, Tripuri, Tulu, Urdu, Vaiphei, Wagdi, Wancho, Yerava, Yerukala/Yerukula, Yimchungre, Zemi, Zou.

There are 5,332 sentences/phrases in Zeliang which are systematically structured based on 159 grammatical categories. Zeliang has 33883 word count and 160654 character count. The corpus is aligned with 146 mother tongues and English, comprising 4,404,845 words (over 4.4 million tokens) and 23,374,289 characters (approximately 23.3 million).

1.6 REFERENCE

- [1] Rejitha K. S. and Narayan Kumar Choudhary. (ed.). (2025). LDC-IL Corpus Insights. CIIL, Mysore. 978-93- 48633-33-0.