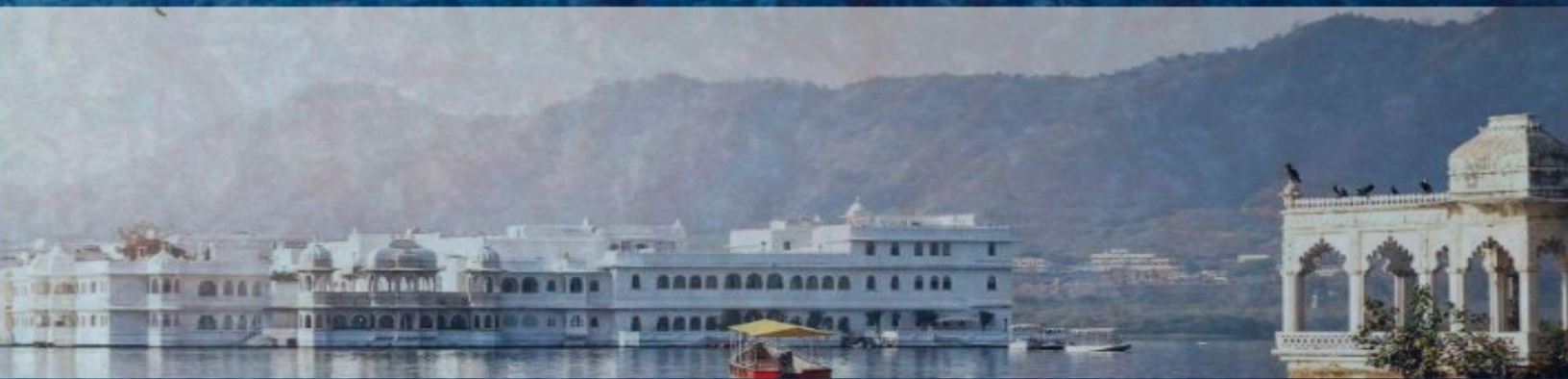
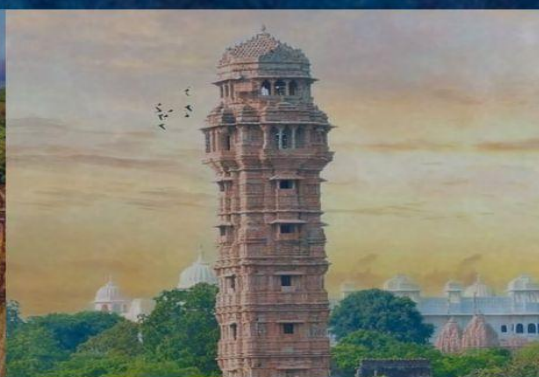
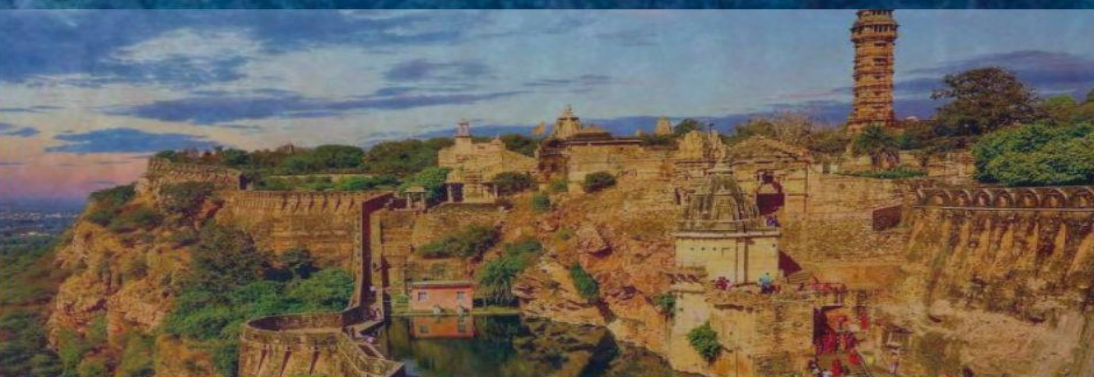




Mewari Parallel Text Corpus: Linguistic Features and Structures



MEWARI 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:

Ankita Tiwari

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

Mewari Parallel Text Corpus: Linguistic Features and Structures

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

Contributors: Dr. Reena Menariya and Dr. Ripudaman Singh Ujjwal

e-ISBN: 978-81-69099-55-4

CIIL Publication No.:1601

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: Ankita Tiwari

TABLE OF CONTENTS

MEWARI PARALLEL TEXT CORPUS.....	5
1.1 LDC-IL PARALLEL TEXT CORPUS.....	5
1.2 CHALLENGES IN CREATING A PARALLEL CORPUS.....	5
1.3 MEWARI: A BRIEF INTRODUCTION.....	6
1.4 MEWARI: A BRIEF LINGUISTIC OVERVIEW	8
1.5 SUMMARY OF MEWARI PARALLEL CORPUS.....	8
1.6 REFERENCE.....	10

MEWARI PARALLEL TEXT CORPUS

Ankita Tiwari

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 MEWARI: A BRIEF INTRODUCTION

Mewari (ISO 639-3: mtr)¹ is a Western Indo-Aryan language spoken primarily in the Mewar region of south-eastern Rajasthan. It functions as the first language of the local population, while Hindi is widely used as a second language. The principal linguistic area encompasses the districts of Udaipur, Rajsamand, Chittorgarh, Bhilwara, and Pratapgarh, with smaller speaker communities in adjoining districts such as Mandsaur and Neemuch in Madhya Pradesh. Among these regions, Chittorgarh is often noted for preserving relatively conservative or “pure” forms of Mewari (Koshy, Abraham, & Raj 2012).

Mewari developed from the Apabhramsha stage of Middle Indo-Aryan, particularly from the Gurjar or Shauraseni Apabhramsha varieties prevalent in north-western India between the 10th and 12th centuries. It belongs to the Indo-Aryan branch of the Indo-European language family and is classified within the Rajasthani subgroup. More specifically, it is regarded as part of the Marwari macro language continuum, representing an eastern variety with close structural and lexical affinities to related speech forms of western India. The Mewar region constitutes the linguistic heartland of Mewari, where it continues to serve as a primary medium of daily communication.²

¹ <https://iso639-3.sil.org/code/mtr>

² https://grokipedia.com/page/Mewari_language

According to the 2011 Census of India, approximately 4,212,262 individuals reported Mewari as their mother tongue.³ Mewari is written using the Devanagari script, which serves as its standard orthographic system in contemporary usage.

1.4 MEWARI: A BRIEF LINGUISTIC OVERVIEW

Mewari forms words primarily through suffixation, which is used for both inflectional and derivational purposes. In addition, compounding and reduplication contribute to vocabulary expansion and semantic emphasis. These processes reflect agglutinative tendencies typical of Western Indo-Aryan languages, allowing grammatical relations to be marked clearly and new meanings to be derived from existing roots. Linguistically, Mewari is characterized by postpositional case marking and relatively conservative preservation of oblique forms compared to Hindi.

The nominal system of Mewari distinguishes two grammatical genders- masculine and feminine and two numbers- singular and plural. Mewari distinguishes three primary cases: direct (used for nominative and accusative functions), oblique (used with postpositions to mark genitive, locative, instrumental, and other relations), and vocative (used for address).

Mewari follows a head-final Subject–Object–Verb (SOV) structure. Postpositions follow noun phrases, and modifiers precede the words they describe. Subordinate clauses generally occur before main clauses. Interrogatives are formed through intonation or the use of the particle *ki*, while negation is expressed by preverbal particles such as *nə-* or *nahi*, without altering the basic alignment system.⁴

1.5 SUMMARY OF MEWARI PARALLEL CORPUS

The Mewari Parallel Text Corpus has been developed using Hindi as the source language. This Corpus connected with English and 146 mother tongues of India. Those mother tongues are Anal, Angami, Apatani, Are, Assamese, Awadhi, Bagri Rajasthani, Bagri, Bagheli/Baghel Khandi, Balti, Bengali, Bhadrawahi, Bharmauri/Gaddi, Bhojpuri, Bilaspuri Kahluri, Bodo,

³ <https://censusindia.gov.in/nada/index.php/catalog/42561>

⁴ https://grokipedia.com/page/Mewari_language

Brajbhasha, Bundeli/Bundel khandi, Chakru/Chokri, Chambeali/Chamrali, Chang, Chhattisgarhi, Chirr, 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, 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, Zeliang, Zemi, Zou.

The Mewari dataset comprises 5,332 sentences/phrases, systematically structured according to 159 grammatical categories. The Mewari portion of the corpus contains 31,252 words and a total of 1,39,194 characters. Moreover, the corpus is aligned with 146 Indian mother tongues in addition to English. Altogether, the multilingual parallel corpus encompasses 4,404,845 words (over 4.4 million tokens) and 23,374,289 characters (approximately 23.3 million).

1.6 REFERENCE

- [1] Iman Sharma. Language Endangerment of Mewari: Study on the Changes in Speakers Attitude
- [2] <https://censusindia.gov.in/nada/index.php/catalog/42561>
- [3] <https://iso639-3.sil.org/code/mtr>
- [4] https://grokipedia.com/page/Mewari_language
- [5] Rejitha K. S. and Narayan Kumar Choudhary. (ed.). (2025). LDC-IL Corpus Insights. CIIL, Mysore. 978-93- 48633-33-0.