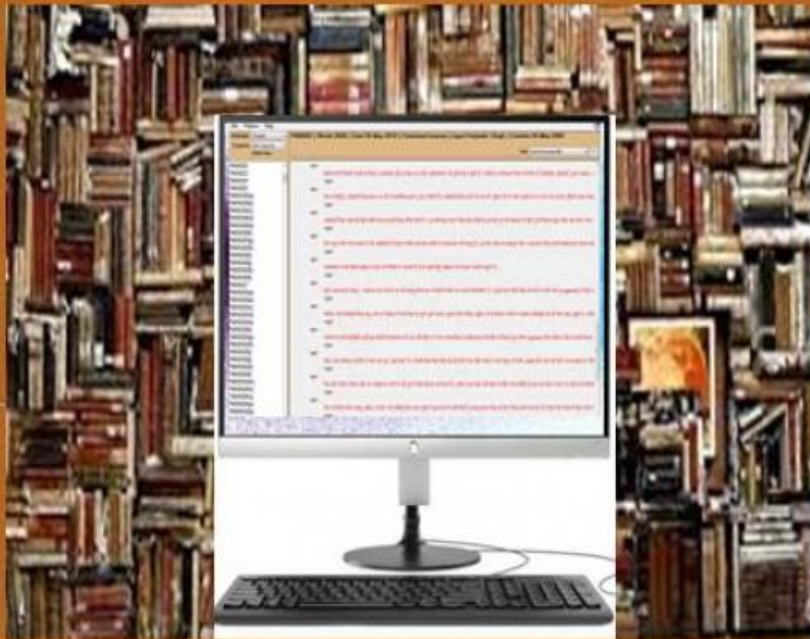


ਉੱਚ-ਮਿਆਰੀ ਪੰਜਾਬੀ ਮੌਲਿਕ ਪਾਠ ਕੋਰਪਸ



**A Gold Standard Punjabi Raw
Text Corpus**



Punjabi Raw Text Corpus

Authors:
Poonam Dhillon
Rajesh N.
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**

Punjabi Raw Text Corpus

This Documentation is a part of LDC-IL's "A Gold Standard Punjabi Raw Text Corpus".

For Further contacts: oic@ldcil.org

All rights reserved

© Central Institute of Indian Languages, Mysore-2018

Authors:

Poonam Dhillon

Rajesh N.

Manasa G.

Dr. Narayan Choudhary

Dr. L. Ramamoorthy

Last updated: December 12, 2018

ISBN 978-81-7343-262-0 Electronic publication of Corpus

Publisher:

Linguistic Data Consortium for Indian Languages

Central Institute of Indian Languages

Mysore, India-570006

Contents

1	LDC-IL Raw Text Corpora; An Overview	1
2	LDC-IL Approach of Sampling.....	3
2.1	Sampling Approach for Books.....	3
2.2	Sampling Approach for Magazines	4
2.3	Sampling Approach for newspaper.....	4
3	LDC-IL Text Corpus categorization	5
3.1	Aesthetics.....	5
3.2	Commerce.....	5
3.3	Mass Media.....	6
3.4	Official Document	6
3.5	Science and Technology.....	6
3.6	Social Sciences	7
4	LDC-IL Text Data encoding and format	7
5	LDC-IL Text Corpus Metadata	8
6	LDC-IL Text Corpus and Naming conventions.....	10
7	Proof Reading.....	11
8	Copyright.....	12
9	Punjabi Raw Text Corpus	13
9.1	Introduction	13
9.2	Peculiarities of Punjabi text	13
9.2.1	Doubling of consonants	14
9.3	Principles of Data Sampling	14
9.4	Fieldworks Undertaken.....	14
9.5	Data Inputting	15
9.6	Validation and Normalization Workshops	15
9.7	Proofreading	15
9.8	Data Extracted from Websites.....	15
9.9	Copyright Consents	15
10	Transliterations in LDC-IL Punjabi text corpus	16

11	Overview of Represented Domains	17
11.1	Aesthetics.....	18
11.2	Commerce.....	18
11.3	Social Science	19
11.4	Mass Media.....	19
11.5	Science & Technology	20

Table of Figures

Table 1: Subcategories of the Category Aesthetics	5
Table 2: Subcategories of the Category Commerce	5
Table 3: Subcategories of the Category Mass Media	6
Table 4: Subcategories of the Category Official Documents	6
Table 5: Subcategories of the Category Science and Technology	7
Table 6: Subcategories of the Category Social Sciences	7
Table 7: Metadata Legends for LDC-IL Text Data.....	8
Table 8: LDC-IL notations for Indian Scheduled languages.....	10
Table 9: Overivew of word count and Character count.....	17
Table 10: Representation of the Domains in Punjabi Text Corpus	17
Table 11: Aesthetics category representationin Punjabi Text Corpus.....	18
Table 12: Commerce category representationin Punjabi Text Corpus.....	18
Table 13: Social Science category representationin Punjabi Text Corpus	19
Table 14: Mass Media category representationin Punjabi Text Corpus.....	19
Table 15: Science & Technology category representation in Punjabi Text Corpus	20

1 LDC-IL RAW TEXT CORPORA; AN OVERVIEW

This is a generic documentation of the LDC-IL raw text corpus which applies to all the languages covered in LDC-IL unless otherwise specified. However, this does not give the specifics of a language dataset.

The objective of language technology is to utilize the facilities of computer, to scientifically analyze language for retrieving verifiable proofs about properties of a language that enable the understanding of multi-dimensional nature of a language. Corpus of a language reflects the nature of the language. The larger and the more representative a corpus, the better it shows its nature.

A corpus is a large collection of language manifestation duly representing its aspects, mainly in text or spoken form. In case of sign language it is the collection of signs in visual form. The electronic text corpus is a collection of pieces of language text in electronic form, selected in accordance with the external criteria to represent, as far as possible, a language or language variety as a source of data for linguistic research. Corpora are one of the major resources for language technology. Computers offer advantages like searching, selecting, sorting and formatting, which eases the language studies. Computers can avoid human bias in an analysis, thus making the result more reliable. Corpus serves as the basis for a number of research tasks within the field of Corpus Linguistics. It is the main resource for many modules of various applications like grammar checkers, spell checkers used in word editors etc. Indian languages often pose difficult challenges for developer community in Natural Language Processing/Artificial Intelligence. The technology developers building mass-application tools/products have for long been calling for availability of linguistic data on a large scale. However, the data should be collected, organized and stored in a manner that suits different groups of technology developers.

Over the years, a lot of efforts have been made to develop text corpora in Indian languages and several agencies have made contributed towards this including the government organizations, academic institutions as well as private bodies. However, the constant greed of more and more electronic data as required by the contemporary machine learning oriented technology models have proved that the data is still not sufficient for all the scheduled languages of India.

Linguistic Data Consortium for Indian Languages (LDC-IL) is one of the Government of India initiatives to develop linguistic corpora in Indian languages. Approved as a scheme in 2007 by the Ministry of Human Resource & Development, Government of India, LDC-IL started functioning at Central Institute of Indian Languages (CIIL), Mysore from April 15, 2008 when human resources got recruited for this scheme. The mission statement for this project is to develop ***“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 acquisition¹.”

The text datasets created under the LDC-IL ambit strives to fill the gap and provide more and more of electronic data for the NLP and language technology community such that the Indian languages get a boost and more of IT applications are available in these languages.

¹ Extract from the *Detailed Project Report* of LDC-IL.

2 LDC-IL APPROACH OF SAMPLING

Developing a written text corpus involves various factors like size of corpus, representativeness, quality of the text, determination of target users, selection of time-span, selection of documents etc. The data for the LDC-IL corpus are collected from books of general interest, textbooks, magazines, newspapers and Government documents of the contemporary text. The data is collected in accordance with prior set of criteria and with the convenience of material such as availability, proper format etc.

As a corpus is supposed to be representative of the language, there is no need to collect all the text from a given book. The representativeness of the corpus depends on a range of different kinds of text categories included in the corpus. LDC-IL corpora try to cover a wide range of text categories that could be representative of the language or language variety under consideration. Corpus representativeness and balance is closely associated with sampling.

LDC-IL collected text corpus from different sources. They are mainly books, magazines, and newspapers. The books are from literature and knowledge text books, magazines and newspapers are web crawled, or keyed in text or both. The newspaper and magazines are great resource of words which are hard to find in books because of the scarcity of those domain specific books in Indian languages.

LDC-IL has different Sampling approach over while extracting text from these three sources.

2.1 SAMPLING APPROACH FOR BOOKS

The books were identified so that the representation of different domains can be catered. After identifying the books, the next step is to extract typically 10 pages of text from it. LDC-IL follows a sampling method to collect the pages from a book. For example, if the book has 100+ pages we collect every 10th page and if the book has 200+ pages we collect every 20th page of the book. If the selected page contains pictures, tables etc, then its next or previous page, which may have the text content, will be chosen for the corpus. Even though one may find rare cases where partial or whole book is selected for the corpus, since the total corpus is going to be very large, such rare cases may not have an impact on balance of corpus. While selecting the book, the LDC-IL's motive is to select from wide variety of domains so that corpus can cover large part of vocabulary and should not miss out certain domain specific words.

Other generic principles that have been normally followed in the sampling tasks across languages are as follows:

- Contents containing obnoxious or vulgar texts have been avoided.
- New editions of the old books having a writing style prior to 1990 were not preferred. Rarely we may have text extracts from such books published prior to 1990 to ensure that the writing style is contemporary.
- For all texts containing short stories, sampling has been made by considering the short stories as a single entity and not based on the whole book containing all the short stories

i.e. each page starting with a new short story have been sampled instead of the usual sampling method based on page numbers of the book.

- The data sampling personnel carried the category and sub-category list for ready reference in the field.
- Text extracts containing poems and formulae have been avoided.
- Pages containing diagrams, tables or figures have been avoided.
- Books containing less than 50 pages are not part of sampling.
- Texts having very small font have been enlarged during photocopying to make it look like 10 to 12 font size.
- If the text contains content other than the intended language, those texts have been avoided if the other language content is longer than one sentence.

2.2 SAMPLING APPROACH FOR MAGAZINES

In case of magazine texts are small and from different domains so the whole magazine is to be considered to be included in corpus discarding advertisements, image captions, and tables etc. Magazine corpus usually includes different types of texts like cookery, health, cinema, stories, contemporary articles, etc.

2.3 SAMPLING APPROACH FOR NEWSPAPER

The newspaper corpus is contemporary text in nature. The text may contain political news, editorials, sports news etc. The news data does not have literary flourish. The news stories are on many unfamiliar domains, religious ideas, scientific principles etc. that have to be conveyed to the common people. So, it is expected that the writers would have captured these domains in a simple and meaningful way. Such write-ups have proper usage of vocabulary, correct language structure and effective phraseology. The newspaper articles may use colloquial, non-standard terms or jargons to attract the readers. The words used need to be expressive and represents the feeling and attitude towards the events. To cover such nuance of the language the newspaper are sampled to be part of the text corpus.

The News items of the paper is sampled based on the domains, classifieds, very small news snippets were avoided. Usually much of the newspaper is keyed.

3 LDC-IL TEXT CORPUS CATEGORIZATION

The LDC-IL corpus shows how people naturally use the language and it does not give imaginary, idealized examples. To satisfy this requirements we needed large amount of data otherwise the frequent items will be from some specific vocabulary or a particular style. Quantitative data gives somewhat accurate results of what occurs frequently and what occurs rarely in the language.

Each text source of corpus is different from others in form, function, content and features. This gives room to classify corpora into different categories. LDC-IL maintains a standard list of categories for which the text is to be collected. LDC-IL Identifies six major categories namely ‘Aesthetics’, ‘Commerce’, ‘Mass Media’, ‘Official Document’, ‘Science and Technology’, ‘Social Sciences’. These categories are further classified into 128 minor categories or sub-categories to cover various domains.

3.1 AESTHETICS

The Aesthetics category is one of the largest contributors to the LDC-IL corpus. This category contains sub-domains from Literature and Fine-arts. The text extracts are from literary sources. It is used to capture literature terms. Aesthetics text is collected from books. The text is probably any standard text which is descriptive in nature. It exhibits the language style of a particular period from which the text is taken. It is an extract of creative writing. It is made up of stories based on fiction, essays on various topics etc. These write-ups are mostly self-expressions of the writer. It captures the flow of language of the writer of the literary text.

The subdomains that are identified for mark-up in corpus under the Aesthetics are given below:

Aesthetics				
Fine Arts-Dance	Literary Texts	Literature-Novels	Autobiographies	Folk Tales
Fine Arts-Drawing	Literature-Criticism	Literature-Plays	Biographies	Folklore
Fine Arts-Hobbies	Literature-Diaries	Literature-Poetry	Cinema	Mythology
Fine Arts-Music	Literature-Essays	Literature-Epics	Culture	Photography
Fine Arts-Sculpture	Literature-Letters	Literature-Speeches	Handicrafts	Humour
Fine Arts-Musical Instruments	Literature-Children’s Literature	Literature-Text Books (School)	Literature-Travelogues	Literature-Science Fiction
				Literature-Short Stories

Table 1: Subcategories of the Category Aesthetics

3.2 COMMERCE

The trade is a part of the society. It exists and operates in association with various groups in society such as customer, suppliers, competitors, banks and financial institutions, Government agencies, trade unions. The trade domain has many domain specific words which need to be part of the corpus. The trade related books will bring such texts to the corpus.

The Subdomains that are identified for mark-up in corpus under the Commerce is given below:

Commerce					
Industry	Accountancy	Share Market	Banking	Business	Career and Employment
Management	Finance	Tourism			

Table 2: Subcategories of the Category Commerce

3.3 MASS MEDIA

Media is an integral part of everyday life for many people all over the world, at work and in the home. The text from this domain is contemporary in nature. The text may contain political news, editorials, or sports news. The major source of the Mass Media text category is newspaper; it contains words which are used in day-to-day life. Structurally, the language of mass media contains exposition, argument, description and narration. It includes different types of write up; consists of structures with different patterns, words and styles. All this is written in a language in which everyone can relate and understand. Some of the media prints are in the form of conversation or question answers. This data usually contains an interviewer and an interviewee. They usually consist dialogues. The interviewee may be a celebrity or a renowned personality from cinema, politics etc. The words used in such text are usually more personal and simple.

The Subdomains that are identified for mark-up in corpus under the Mass Media is given below:

Mass Media					
Article	Classifieds	General News	Obituary	SMS	Religious/Spiritual News
Business News	Discussions	Interviews	Political	Social	Sports News
Cinema News	Editorial	Letters	Speeches	Weather	Health

Table 3: Subcategories of the Category Mass Media

3.4 OFFICIAL DOCUMENT

The usage of language in official documents is highly standard, unambiguous, straight forward and structurally modified. The communication intended in official documents are intended about some action, or some enquiry or proceedings of some assemblies. This text usually is to get the due representation of such domain specific terminologies of administration, official document category is included.

The Subdomains that are identified for mark-up in corpus under the Official Document is given below:

Official Document			
Administration	Legislature	Parliamentary/Assembly Debates	Police Documents

Table 4: Subcategories of the Category Official Documents

3.5 SCIENCE AND TECHNOLOGY

The science and technology domain contains text extracts from various scientific books, articles of magazines, journals etc. These texts are also called as knowledge texts. The language structure and usage of words are different from the language of day-to-day life. The terminologies that are from this domain will have highest number of loan words because the subject in the text is usually global. To get the due representation of such domain specific terminologies, the Science and Technology category is included.

The Subdomains that are identified for mark-up in corpus under the Science and Technology is given below:

Science and Technology					
Agriculture	Biotechnology	Engineering-Civil	Forestry	Medicine	Statistics
Architecture	Botany	Engineering-Electrical	Geology	Micro Biology	Astrology
Textile Technology	Educational Psychology	Engineering-Electronics Communication	Text Book (Science)	Computer Sciences	Language Technology
Chemistry	Naturopathy	Engineering-Mechanical	Horticulture	Oceanology	Veterinary
Ayurveda	Criminology	Engineering-Others	Astronomy	Physics	Film Technology
Bio Chemistry	Homeopathy	Environmental Science	Logic	Psychology	
Biology	Yoga	Engineering-Chemical	Mathematics	Sexology	Zoology

Table 5: Subcategories of the Category Science and Technology

3.6 SOCIAL SCIENCES

Language is a medium for creation and maintenance of human society so language in social sciences category correlates the linguistic features of the dynamic society. Human development and reformation happening in different communal context hence all the social knowledge and reality could be reflected in this text category.

The Subdomains that are identified for mark-up in corpus under the Social Sciences is given below:

Social Sciences					
Anthropology	Food and Wellness	Personality Development	Physical Education	Text Book (Social Science)	Philosophy
Archaeology					Journalism
Demography	Fisheries	Library Science	Law	Sports	Geography
Economics					Religion / Spiritual
Education	Home Science	Political Science	Public Administration	Health and Family Welfare	Sociology
Epigraphy					Linguistics

Table 6: Subcategories of the Category Social Sciences

4 LDC-IL TEXT DATA ENCODING AND FORMAT

The collected data should be encoded in a machine readable form for further analysis. While storing the data one has to keep some standards so that the data is easy to store and retrieve in long term. The encoding being used in LDC-IL Text corpus is Unicode and stored in XML format. Large scale language resource depends on the metadata. Metadata is an authentic source to prove the quality of the data. Metadata should have the subject information, source information and encoding information.

The selected text along with metadata information is indexed with a five digit unique number to get keyed-in. Each text fragment of selected book is typed as corpus file with xml extension. The given unique Index number gets prefixed with the LDC-IL notations which make the filename of the XML file. Sometimes the XML file names carry small case alphabets enclosed in braces. This is done if the book title carries different type of textual topics, so that each chapter, in the selected book title which may be related to different topics, chapters etc., can be differentiated. This helps the text content get categories based on the context.

5 LDC-IL TEXT CORPUS METADATA

It is imperative to maintain metadata of the entire data collection for linguistic analysis. The collected data are arranged with its metadata information such as its category, subcategory, title of the text, author name, source, publisher name, year of publication, page numbers etc. This information helps the users to retrieve the data easily from the database/repository. Metadata gives authenticity to the text by way of providing the details of how the data was created in the first instance and what is its content about. The following table shows the legend used in the metadata and provides description of them.

#	Legend	Description
1	Filename	Represented by "docID" tag in the XML files. This is a unique file number across the datasets.
2	Project Description	This gives a brief of the project under which the file was generated. As CIIL has been involved into corpus creation over a long period time, including before the inception of LDC-IL scheme, there might be some data for a few languages which might have come from different projects e.g. the CIIL Corpus or CIIL-KHS corpus. This field indicates the source of the project.
3	Sampling Description	This information is a verifiable proof for the corpus. It will have the information of selected page numbers of the book for corpus.
4	Category	Specifies the domain of the text.
5	Subcategory	Specifies the sub-domain of the text.
6	Text	Specifies the type of the source text i.e. whether its origin is a book, a magazine or a newspaper.
7	Title	Specifies the title of the source text. It contains mostly books but if magazines or newspapers occur, their respective are provided here.
8	Volume	Specifies volume number the title, if any.
9	Issue	Specifies issue number the title, if any.
10	Text Type	Is mostly blank however sometimes it is used to provide the broad topic of the news items e.g. whether it is a political news or editorial or sports news etc.
11	Headline	This information is a verifiable proof for the corpus. This is normally the heading of the chapter of the selected sample. Gives the fine tuned information of the topic present in particular file.
12	Author	Specifies the name of the author.
13	Editor	Specifies the name of the editor.
14	Translator	Specifies the name of the translator.
15	Words	Specifies the total number of words in the file.
16	Letters	Specifies the total number of UTF8 characters in the file.
17	Publishing Place	Specifies the place where the title was published.
18	Publisher	Specifies the name of the publisher.
19	Published Year	Specifies the publishing year.
20	Index	Is the index number or ID of the file. It is noted inside the XML file. It is mostly the same as the file name.
21	Date	Date when the file was digitized/inputted.
22	Input	Name of the Data Inputter, if the file has been typed.
23	Proof	Name of the Proof reader.
24	Language	Name of the language.
25	Script	Name of the script the text is written in.

Table 7: Metadata Legends for LDC-IL Text Data

Typical Metadata Mark-ups in a text corpus file structure is given below.

<?xml version="1.0" ?>			
<?xml-stylesheet type="text/css" href="home.css"?>			
<Doc id="mal-w-media-		ML00172	"
<Header type="text">		lang="Malayalam">	
<encodingDesc>			
<projectDesc>	CIIL-Malayalam Corpora, Monolingual Written Text		</projectDesc>
<samplingDesc>	Simple written text only has been transcribed. Diagrams, pictures and tables have been omitted. Samples taken from page 30-31,50-51,70-71,94-95,114-115,132-133,152-153,172-173,192-193,210-211		</samplingDesc>
</encodingDesc>			
<sourceDesc>			
<biblStruct>			
<source>			
	<category>	Aesthetics	</category>
	<subcategory>	Literature-Novel	</subcategory>
	<text>	Book	</text>
	<title>	Kalapam	</title>
	<vol>		</vol>
	<issue>		</issue>
</source>			
<textDes>			
	<type>		</type>
	<headline>		</headline>
	<author>	ShashiTharoor	</author>
	<editor>		</editor>
	<translator>	Thomas George	</translator>
	<words>	2745	</words>
</textDes>			
<imprint>			
	<pubPlace>	India-Kottayam	</pubPlace>
	<publisher>	DC Books	</publisher>
	<pubDate>	2006	</pubDate>
</imprint>			
<idno type="CIIL code">		Kerala University Campus Library- 13535	</idno>
<index>		ML00172	</index>
</biblStruct>		</sourceDesc>	
<profileDesc>		<creation>	
	<date>	26-Apr-2010	</date>
	<inputter>	Remya K	</inputter>
	<proof>		</proof>
</creation>			
<langUsage>		Malayalam	</langUsage>
<ScriptUsage>		Malayalam	</ScriptUsage>
<wsdUsage>			
<writingSystem id="ISO/IEC 10646"> Universal Multiple-Octet Coded Character Set (UCS). </writingSystem>			
</wsdUsage>			
<textClass>			
<channel mode="w">		Print	</channel>
<domain type="public">			</domain>
</textClass>		</profileDesc> </Header>	
<text> <body>			
<p>			
<p>			
</text> </body> </Doc>			

6 LDC-IL TEXT CORPUS AND NAMING CONVENTIONS

The selected hardcopies were marked for sampling and given to typists by concerned language experts. LDC-IL has built an in-house corpus developing application and stores it in a repository database. The samples get typed in xml format through a software application built for it in LDC-IL. Each sampling is a corpus file and gets typed and saved in Unicode standards. Each corpus file has unique filename. One can say the corpus is indexed through filenames. Typically each corpus file is an extract of a book of a particular title. The LDC-IL corpus file name follows certain naming convention. The naming convention is based on language and source of text. Every scheduled language has a notation for each kind of source of corpus. The notation is prefixed to a five digit number to create a unique corpus filename.

The LDC-IL notations for Indian Scheduled languages are given below.

#	Language	ISO 639 Language Code	Script	Notation as per Source of Corpus			
				Book	Magazine	News Paper	News Web
1	Assamese	asm	Assamese	AS	ASM	ASN	ASNW
2	Bengali	ben	Bengali	BE	BEM	BEN	BENW
3	Bodo	brx	Devanagari	BD	BDM	BDN	BDNW
4	Dogri	doi	Devanagari	DG	DGM	DGN	DGNW
5	Gujarati	guj	Gujarati	GJ	GJM	GJN	GJNW
6	Hindi	hin	Devanagari	HN	HNM	HNN	HNNW
7	Kannada	kan	Kannada	KA	KAM	KAN	KANW
8	Kashmiri	kas	Persio-Arabic	KS	KSM	KSN	KSNW
9	Konkani	kok	Devanagari	KO	KOM	KON	KONW
10	Maithili	mai	Devanagari	MT	MTM	MTN	MTNW
11	Malayalam	mal	Malayalam	ML	MLM	MLN	MLNW
12	Manipuri	mni	Bengali	MN	MNM	MNN	MNNW
13	Marathi	mar	Devanagari	MA	MAM	MAN	MANW
14	Nepali	nep	Devanagari	NP	NPM	NPN	NPNW
15	Odia	ori	Odia	OD	ODM	ODN	ODNW
16	Punjabi	pan	Gurmukhi	PN	PNM	PNN	PNNW
17	Sanskrit	san	Any Script	SA	SAM	SAN	SANW
18	Santali	sat	OlChiki	SN	SNM	SNN	SNNW
19	Sindhi	snd	Persio-Arabic	SI	SIM	SIN	SINW
20	Tamil	tam	Tamil	TA	TAM	TAN	TANW
21	Telugu	tel	Telugu	TE	TEM	TEN	TENW
22	Urdu	urd	Persio-Arabic	UR	URM	URN	URNW

Table 8: LDC-IL notations for Indian Scheduled languages

Consider the example of Malayalam, The text taken from Malayalam book for LDC-IL Malayalam Text Corpus always starts with ‘ML’ followed by 5 digit numbers which is continuous, where as text collected from Malayalam Magazine starts with ‘MLM’ followed by 5 digit numbers. If the source is from Newspaper then ‘MLN’ notation will be followed where as if the News is taken from Web source ‘MLNW’ will be used as notation.

In certain cases, if the book is chaptered, the headline of each chapter changes, to capture the change of the topic. If the language experts wish to break the sampling of a book into different smaller files, then the filename will get attached with roman small letter suffixed and enclosed in braces.

Such filenames could be ‘ML00001 (a)’, ‘ML00001 (b)’, ‘ML00001 (c)’, ‘ML00001 (d)’ etc.

7 PROOF READING

Once it is in digital form, the same is proofread so that it is free from any kind of typographical errors. Proofing is the next process of corpus building. Since the typed corpus may carry errors because of various reasons like speed of the typist and typist not belonging to the language community, the proofing is done by the language experts.

While proofing of a corpus file is done in LDC-IL, the following things are taken care of

1. Removing the poetic text, if any poem or poetic structure occurs within the running text
2. If there are incomplete sentences typed (generally at the end of the paragraph) the sentence is removed up to the logical ending of the previous sentence.
3. Verifying the difference between the visargaha and colon ‘ : ’ symbol, and to ensure that the correct symbol/punctuation is used in the correct place.
4. During Content cleaning focus stays on the corrections of typographical errors and spacing. If there is a space preceding a punctuation mark, space is removed, unless it is there in the actual text itself (i.e. hard copy of the text).
5. If there is any mismatch between the hard copy and the input corpus file, it is ensured that the corpus file should be faithful to hard copy.
6. It is ensured that the Title, Author, Headline fields of the XML files is written in Roman using the LDC-IL transliteration scheme. The LDC-IL Transliteration scheme can be referred on the LDC-IL website. Also, the LDC-IL transliteration tool from Roman to Indian Scripts and vice versa is available for download on the LDC-IL website.

Link to download LDC-IL Transliteration Scheme:

<http://ldcil.org/Tools/CorporaToolsPackage/LDC-IL%20Transliteration%20Scheme.pdf>

Link to download the LDC-IL Transliteration Tool (.exe file):

<http://ldcil.org/Tools/LDC-IL%20Transliterator.zip>

Proof reading is used to correct clear cases of spelling mistakes, splitting sentences or words, removing unnecessary repeated paragraphs, sentences, phrases, words. Moreover, it includes removing unwanted texts from the corpus such as foreign script sentences and incorrect use of ungrammatical sentences.

8 COPYRIGHT

Anyone intending to put together a corpus for commercial purposes must always obtain the permission from the publishers of the source texts. Many commercially available corpora contain texts from a large number of sources and obtaining permission to use these can be a very cumbersome and financially costly process. However, LDC-IL took up the task and managed to get the consent of most of the copyright holders or has at least communicated to them that the text extracts from their sources are being used in the language sampling task which may also be used commercially.

Considering LDC-IL is a government initiative taken up in the larger public interest and the corpus is used for the development of language, most of the publishers and authors generously agreed to archive the samples of their text materials in corpus. Some of the authors even suggested and offered their other content which are not yet part of the LDC-IL corpus. Government publishers too expressed no objections regarding since LDC-IL itself is an initiative of Govt. of India. Private publishers also gave permission considering that LDC-IL is only using a part of a text, and it will not harm their business anyway. LDC-IL thanks all of them for the co-operation.

For some of the content where we have not yet got the explicit consent of the copyright holders, we have sent them the letters asking for the same. If any of the copyright holders disagree to consent, they may write so to us and their respective text will be removed from the sampling corpus and the same will be intimated to all the license holders of the respective dataset and they will have to abide by it.

9 PUNJABI RAW TEXT CORPUS

9.1 INTRODUCTION

Punjabi is the principal and administrative language of Punjab. Punjabi is a tonal language with three tones: high falling, low rising and level. Punjabi is not only spoken in Punjab in India it is also a language of Lehnda Punjab in Pakistan. In Pakistani Punjabi is the second most widely-spoken language in Pakistan but has no official status. Punjabi is an Indo-Aryan language. It is derived from Sanskrit through Prakrit languages and later *Apabhraṃs*. There was no such form of Punjabi language in the beginning that we see today. With the flow of time, it has emerged in the present form. This same Punjabi language is being written in two epigraphs in Gurmukhi and Shahmukhi script. In our Eastern Punjab it is being used in Gurmukhi and Lehnda Punjab (Pakistan) using Shahmukhi script.

Punjabi language with Gurmukhi script:

In 16th century Guru Angad Dev Ji, the second Sikh guru was standardised the Gurmukhi alphabet from the Landa alphabet. Gurmukhi is written from right to left. This script has 10 vowels and 29 consonants and 5 Perso-arabic consonants. It also has two semivowels / y / and / v /.

Punjabi language with Shahmukhi script:

The Shahmukhi alphabet is a version of Perso-Arabicalphabet and used to write Punjabi in Pakistan. Shahmukhi is written from left to right. This script has 10 vowels and 46 consonants and 10 mixed words.

Punjabi is written in Shahmukhi scripts as well. ‘Shahmukhi’ is a variant of ‘Perso-Arabic’ script. LDC-IL Punjabitext corpus is collected in Gurmukhi script of contemporary usage.

Punjabi text corpus is collected from various libraries in Punjab mostly from Patiala. The greater part of the text has been taken from NRLC library and Punjabi University Patiala library. LDC-IL tried to cover the entire category in its standard list. Some categories like novel, short stories has huge amount of books but some categories like physics, chemistry, economics have very less amount of books. Literary texts are easily available in Punjabi but getting scientific text is very difficult. Some categories like epigraphy, finance, oceanology text are not available in Punjabi.

9.2 PECULIARITIES OF PUNJABI TEXT

The Corpus of Punjabi text can be broadly classified into two: literary text and non-literary text. These two explicitly show their differences in terms of frequency of word usage and variety that it brings into corpus. Literary texts are texts that are narrative and it contains elements of fiction. Novels, short stories, plays are examples of literary text. Non-literary texts are texts whose primary purpose is to convey information. Examples of non-literary texts are text about various scientific or technical subjects, legal documents, articles in academic journals. In literary text, language has emotional elements, cultural information, dialectical variations, ambiguity etc. But technical or scientific terms, foreign words etc. have widely appeared in non-literary texts.

9.2.1 Doubling of consonants

In Punjabi Text corpus the ‘ੌ’ (GURMUKHI TIPPI) is used for the nasalization rather than ‘ੌ’ (GURMUKHI SIGN ADAK BINDI).

The other predominant feature in Gurumukhi Script is the usage of ‘ੌ’ (GURMUKHI ADDAK) which doubles following consonant to which it attaches. Unlike other Indian languages which uses Virama (Halanth) to make the Half-letter followed by full letter, Gurmukhi uses ADDAK.

While processing text this ADDAK has no value until the next consonant is known. This will create problem in text processing applications like transliterator, character based morph analyser etc. That are analysing and processing character by character of the given text. The text processors need to be enabled with extra feature of checking immediate next character of the ADDAK.

The LDC-IL Corpus uses ADDAK as it occurs naturally in Punjabi Text written in Gurumukhi Script.

9.3 PRINCIPLES OF DATA SAMPLING

Punjabi text data sampling strictly followed the generic guidelines of LDC-IL text corpus collection which are noted in the generic LDC-IL corpus documentation.

9.4 FIELDWORKS UNDERTAKEN

Punjabi text corpus is collected from various libraries in Punjab, mostly from Patiala. The text materials were collected by conducting three fieldworks undertaken in the period from 2008 to 2010. The greater part of the text has been taken from NRLC, Patiala and Punjabi University Patiala library. Overall, the following libraries served as the source of the Punjabi text corpus:

- Northern Regional Language Centre, Patiala
- Punjabi University, Patiala
- Khalsa College, Patiala
- Guru Nanak Dev University, Amritsar
- Regional Campus, Jalandhar

Mostly collected text materials have been published from Punjab and New Delhi.

An attempt has been made to cover the entire category in its standard list. Some categories like criticism, novel and short stories have huge amount of books but some categories like physics, chemistry, economics have very less amount of books. Literary texts are easily available in Punjabi but getting scientific text is very difficult. Some categories like epigraphy, finance, oceanology text are not available in Punjabi.

Collecting text data from the field is a difficult job. Most of the libraries do not allow to take huge amount of text from their shelves at a time, because it is against their rules and principles. For a particular period, they issue maximum three or four books. Even if the librarian allowed to take many books at a time, the photocopy kiosk had issues as there was a long queue.

Sometime Xerox attendants refused to photocopy randomly selected pages because of the long queue waiting and it takes up more time for them to turn the pages compared to continuous page photocopying they are accustomed to. It was another issue that the field worker/linguist had to carry a huge list of photocopy bundles with them which was many times cumbersome to travel with.

Despite all the issues as above, the linguists working on the data collection had to deal with and get going.

9.5 DATA INPUTTING

All the text has been typed in Unicode using the In Script Keyboard directly onto the XML files. The data has been inputted by Harjinder Singh, Gurmeet Kaur, Harpreet Kaur, Kulwant Singh, native speaker of Punjabi, but Radhika M, Syeda Aliya Habeeba native speaker of Kannada.

9.6 VALIDATION AND NORMALIZATION WORKSHOPS

A 5-day workshop was conducted at Linguistic Data Consortium from November 28 to December 1, 2011 with Prof. Joga Singh Department of Linguistics, Punjabi University, Patiala, Prof. Baldev Singh Cheema Department of Punjabi, Punjabi University, Patiala and Prof. Sukhwinder Singh Sangha from Department of Punjabi, Regional Campus, Jalandhar as experts. The experts suggested that the Punjabi text corpus should remain true to the text.

9.7 PROOFREADING

Punjabi text data has been proof read by internal and external resource persons. The text has always been kept true to the printed material and typos, if any, occurring at the time of typing have only been corrected. Some text cleaning workshops were conducted using external resources wherein the Punjabi text was cleaned/proofread by the native speakers. An account of such workshops is as below:

1. July 2010
2. 24th Dec. 2012 – 28th Feb. 2013
3. 05 Aug-18 Sept 2013
4. 03 Oct-07 Dec 2016
5. 01 May-05 July 2018

The printed materials collected for the corpus is contemporary, mainly published after 1980.

9.8 DATA EXTRACTED FROM WEBSITES

Punjabi News corpus data is extracted from News websites of "Ajit Weekly" (www.ajitweekly.com)," Charhdikala" (www.charhdikala.com), " Europe Vich Punjabi" (www.europevichpunjabi.com), " Pardes News " (www.pardesnews.com)" Parvasi " (www.parvasi.com) " Punjab Express " (www.punjabexpress.com)," Punjabi Webdunia " (www.punjabi.webdunia.com), and" Quami Ekta " (www.quamiekta.com). The news content was categorized based on the content of the text and archived. The period of selection of the news corpus ranges from 2008 to 2010.

9.9 COPYRIGHT CONSENTS

The Punjabi text corpus have been collected from various sources and the copyright for the same stays with different sources. However, for the purposes of this corpus, consents have been sought from all the stakeholders. Most of the copyrights (around 73%) belong to private parties with only 27% belonging to the government agencies, either state or the central.

10 TRANSLITERATIONS IN LDC-IL PUNJABI TEXT CORPUS

For easy reference and uniformity of metadata, some entries in the metadata file, namely ‘Title’, ‘Headline’, ‘Author’, ‘Editor’, ‘Translator’ are transliterated from Gurmukhi to Roman letters. Numeric characters were transliterated from Gurmukhi to Hindu-Arabic system.

For such purpose the LDC-IL transliteration scheme for Gurmukhi to Roman characters is given below:

LDC-IL Transliteration Schema										
Gurmukhi characters to Roman and Gurmukhi Numerals to Hindu-Arabic										
Vowels										
ਅ	ਆ	ਇ	ਈ	ਉ	ਊ	ਏ	ਐ	ਓ	ਔ	
	ਾ	ਿ	ੀ	ੁ	ੂ	ੇ	ੈ	ੋ	ੌ	
a	A	i	I	u	U	E	ai	O	au	

Consonants				
ਕ	ਖ	ਗ	ਘ	ਙ
ka	kha	ga	gha	ng'a
ਚ	ਛ	ਜ	ਝ	ਞ
ca	cha	ja	jha	nj'a
ਟ	ਠ	ਡ	ਢ	ਣ
Ta	Tha	Da	Dha	Na
ਤ	ਥ	ਦ	ਧ	ਨ
ta	tha	da	dha	na
ਪ	ਫ	ਬ	ਭ	ਮ
pa	pha	ba	bha	ma

Symbols			
ੱ	ੰ	ਂ	ਃ
Null	m'	M	H

ਯ	ਰ	ਲ	ਵ	ੜ	ਸ਼	ਖ਼	ਗ਼	ਜ਼	ਫ਼	ਲ਼
ya	ra	la	va	Ra	sha	Kh'a	g'a a	j'a	ph' a	La

Numerals(Punjabi to Hindu-Arabic)									
੦	੧	੨	੩	੪	੫	੬	੭	੮	੯
0	1	2	3	4	5	6	7	8	9

The greyed out characters are obsolete. They may rarely present in the current LDC-IL corpus.

11 OVERVIEW OF REPRESENTED DOMAINS

LDC-IL Punjabi Text Corpus size is: 1,01,25,770 words and characters count is 5,08,24,349 drawn from 2,470 different titles, including the extracts from newspapers and magazines. The data can be categorized into two classes of typed+cleaned and crawled. The crawled data has been crawled mainly from news websites and archived using the standard processing of LDC-IL text corpus preparation.

The following table gives a summary of the typed and crawled text of the Punjabi Raw Text Corpus.

Text Type	Word Count	Keystroke/Character Count
Typed+Cleaned	97,55,905	4,89,97,317
Crawled	3,69,865	18,27,032
Total	1,01,25,770	5,08,24,349

Table 9: Overview of word count and Character count

The representation of the five major domains covered has been shown in the table below:

Domain	Domain Word Count	Percentage
Aesthetics	41,90,199	41.38%
Commerce	56,205	00.56%
Social Sciences	12,20,366	12.05%
Mass Media	42,74,922	42.22%
Science & Technology	3,84,078	03.79%
Total	1,01,25,770	100.00%

Table 10: Representation of the Domains in Punjabi Text Corpus

As each domain has several sub-domains, the following table shows the representation of the several domains, both within the domain and across all the domains.

11.1 AESTHETICS

The Aesthetics category of Punjabi text corpus covers 22 sub-categories bearing a total of 41,90,199 words along with the overall percentage of 41.38%

%. The representational details are given in the table below.

#	Sub Category	Word Count	% within Subdomain	Overall Percentage
1	Autobiographies	1,45,184	3.46%	1.43%
2	Biographies	2,60,595	6.22%	2.57%
3	Cinema	73,865	1.76%	0.73%
4	Culture	1,37,921	3.29%	1.36%
5	Fine Arts-Dance	27,289	0.65%	0.27%
6	Fine Arts-Drawing	8,780	0.21%	0.09%
7	Fine Arts-Sculpture	36,964	0.88%	0.37%
8	Fine Arts-Music	35,652	0.85%	0.35%
9	Folklore	1,04,858	2.50%	1.04%
10	Humor	1,584	0.04%	0.02%
11	Literary Texts	4,27,152	10.19%	4.22%
12	Literature-Children's Literature	1,693	0.04%	0.02%
13	Literature-Criticism	14,83,799	35.41%	14.65%
14	Literature-Diaries	25,845	0.62%	0.26%
15	Literature-Letters	10,501	0.25%	0.10%
16	Literature-Novels	4,71,785	11.26%	4.66%
17	Literature-Plays	54,775	1.31%	0.54%
18	Literature-Short Stories	6,65,293	15.88%	6.57%
19	Literature-Speeches	83,710	2.00%	0.83%
20	Literature-Travelogues	1,07,341	2.56%	1.06%
21	Literature-Text Books (Schools)	23,192	0.55%	0.23%
22	Mythology	2,421	0.06%	0.02%
	Total	41,90,199	100%	41.38%

Table 11: Aesthetics category representation in Punjabi Text Corpus

11.2 COMMERCE

The Commerce category of Punjabi text corpus covers 2 sub-categories bearing a total of 56,205 words along with the overall percentage of 0.56%. The representational details are given in the table below.

#	Sub Category	Word Count	% within Subdomain	Overall Percentage
1	Business	45,159	80.35%	0.45%
2	Management	11,046	19.65%	0.11%
	Total	56,205	100%	0.56%

Table 1112: Commerce category representation in Punjabi Text Corpus

11.3 SOCIAL SCIENCE

The Social Science category of Punjabi text corpus covers 17 sub-categories bearing a total of 12,20,366 words along with the overall percentage of 12.05%. The representational details are given in the table below.

#	Sub Category	Word Count	% within Subdomain	Overall Percentage
1	Economics	1,58,038	12.95%	1.56%
2	Education	86,299	7.07%	0.85%
3	Food and Wellness	1,261	0.10%	0.01%
4	Geography	18,317	1.50%	0.18%
5	Health and Family Welfare	16,859	1.38%	0.17%
6	History	1,69,151	13.86%	1.67%
7	Home Science	7,115	0.58%	0.07%
8	Journalism	46,378	3.80%	0.46%
9	Law	22,057	1.81%	0.22%
10	Library Science	20,977	1.72%	0.21%
11	Linguistics	1,33,404	10.93%	1.32%
12	Physical Education	54,906	4.50%	0.54%
13	Political Science	1,25,997	10.32%	1.24%
14	Public Administration	76,345	6.26%	0.75%
15	Religion/Spiritual	1,81,472	14.87%	1.79%
16	Sociology	91,775	7.52%	0.91%
17	Sports	10,015	0.82%	0.10%
	Total	12,20,366	100 %	12.05%

Table 13: Social Science category representation in Punjabi Text Corpus

11.4 MASS MEDIA

The Mass Media category of Punjabi text corpus covers 14 sub-categories bearing a total of 42,74,922 words along with the overall percentage of 42.22%. The representational details are given in the table below.

#	Sub Category	Word Count	% within Subdomain	Overall Percentage
1	Business News	1,20,017	2.81%	1.19%
2	Cinema News	1,08,318	2.53%	1.07%
3	Classifieds	785	0.02%	0.01%
4	Discussions	1,610	0.04%	0.02%
5	Editorial	9,14,446	21.39%	9.03%
6	General News	22,56,487	52.78%	22.28%
7	Health	14,578	0.34%	0.14%
8	Interviews	21,589	0.51%	0.21%
9	Letters	3,094	0.07%	0.03%
10	Political	4,82,448	11.29%	4.76%
11	Religious/Spiritual News	26,957	0.63%	0.27%
12	Social	27,810	0.65%	0.27%
13	Speeches	3,054	0.07%	0.03%
14	Sports News	2,93,729	6.87%	2.90%
	Total	42,74,922	100%	42.22%

Table 14: Mass Media category representation in Punjabi Text Corpus

11.5 SCIENCE & TECHNOLOGY

The Social Science category of Punjabi text corpus covers 17 sub-categories bearing a total of 3,84,078 words along with the overall percentage of 3.79%. The representational details are given in the table below.

#	Sub Category	Word Count	% within Subdomain	Overall Percentage
1	Agriculture	42,294	11.01%	0.42%
2	Astrology	11,990	3.12%	0.12%
3	Ayurveda	40,680	10.59%	0.40%
4	Bio Chemistry	24009	6.25%	0.24%
5	Botany	21,913	5.71%	0.22%
6	Computer Sciences	44,164	11.50%	0.44%
7	Criminology	6,175	1.61%	0.06%
8	Environmental Science	22,797	5.94%	0.23%
9	Forestry	9,448	2.46%	0.09%
10	Homeopathy	29,850	7.77%	0.29%
11	Medicine	31,978	8.33%	0.32%
12	Naturopathy	6,199	1.61%	0.06%
13	Physics	19,661	5.12%	0.19%
14	Psychology	21,398	5.57%	0.21%
15	Text Book (Science)	14,584	3.80%	0.14%
16	Yoga	9,903	2.58%	0.10%
17	Zoology	27,035	7.04%	0.27%
	Total	3,84,078	100%	3.79%

Table 15: Science & Technology category representation in Punjabi Text Corpus