

Development and Integration of an Odia Stemmer in Dspace for Odia Search Engine



Gouranga Charan Jena, Siddharth Swarup Rautaray

Abstract: Stemmer is used for reducing inflectional or derived word to its stem. This technique involves removing the suffix or prefix affixed in a word. It can be used for information retrieval system to refine the overall execution of the retrieval process. This process is not equivalent to morphological analysis. This process only finds the stem of a word. This technique decreases the number of terms in information retrieval system. There are various techniques exists for stemming. Here a new hybrid stemmer has developed named as "Mula" for Odia Language. It is a combination of brute force and enhanced suffix stripping approach for Odia language. The new born stemmer is both computationally inexpensive and domain independent. We have integrated this stemmer in existing Dspace for Odia text retrieval System. The results are commendable and suggest that the new stemmer can be used effectively in Odia Search Engine. The proposed stemmer also handles over-stemming and under-stemming effectively.

Keywords: Derivational Suffixes, inflectional words, Odia Stemmer, Information Retrieval, Brute force, DSpace etc.

I. INTRODUCTION

Stemming is a technique for removing inflectional or derived word to its stem. This technique is used to remove the suffix or prefix affixed in a word. This process finds stem of a word. Stemmer is essential for information retrieval system to refine the performance of the system. It technique is not equivalent to morphological analysis. Its primary objective is to decrease the number of terms in an information retrieval system. Stemming technique can be used in information retrieval to decrease as many related words to a common form that is not in base form. For example, the English word "Computation" has different inflections such as 'Comput', 'Compute', 'Computing', 'Computes' etc. In this case stemmer can be used to reduce derived words into its root or stem word. Many stemmers had been developed for different languages, which reduce a word to its root/stem form. It ultimately reduces the index file size in an information retrieval system. In this way we can improve recall (i.e. the number of documents retrieved in response to a query.) of an IR system by effectively using stemmer in the background. Since many derivational words are mapping into one word i.e. root or stem. It ultimately reduces the volume of the index files in the IR system. There are several types of stemming algorithms exists and it differs in respect to their performance and accuracy.

Revised Manuscript Received on February 28, 2020.

* Correspondence Author

Gouranga Charan Jena*, Anand Pani School of Computer Engineering, KIIT DU, Odisha, India. Email- jenagouranga2000@gmail.com

Siddharth Swarup Rautaray, School of Computer Engineering, KIIT DU, Odisha, India. Email- siddharthfcs@kiit.ac.in

© The Authors. Published by Blue Eyes Intelligence Engineering and Sciences Publication (BEIESP). This is an open access article under the CC-BY-NC-ND license <http://creativecommons.org/licenses/by-nc-nd/4.0/>

There are various algorithms used to find a stem of a word i.e. (a) Brute-force algorithms: It uses a lookup table that contains derived words with their corresponding roots.

To find the root/stem of a word, the table is queried to find a matching inflection word. If a matching inflection word is found, then corresponding root returned. Otherwise, it fails. (b) Suffix-stripping Approach: It does not depend on any lookup table that consists of derived or inflected words and their root word relations. It simply uses a set "rules" which drives the algorithm. It finds the root/stem of the given input word based on that rules. (c) Lemmatization algorithms: This technique also called as text normalization. In Lemmatization root word is called Lemma. The POS is first identified of that language and then an attempt will be made to find the stem. The stemming rules will change based on a word's POS of that language. Lemmatization process ensures that the root or stem of the inflected words belongs to that language (d) Stochastic algorithms: It is based on probability method to detect the root form of a word. This trained on a table of root words to inflected words relations to develop a probabilistic model. (d) Affix Removal Approach: The name clearly suggests this approach is related to removing the suffixes or prefixes of a word. An Affix may be a prefix or a suffix. It comes under truncating method of the stemming algorithm. We found affixes are connected with nouns in Odia language. One can opt any of the above technique while designing stemmer. (e) Hybrid approach: This technique combines more than two methods as discussed above. It may merge the rule-based technique along with the probability method. (f) N-Gram Modeling: Many stemming methods used in the n-gram technique of a word to select the correct stem for a word. Stemming plays a vital role to handle the vocabulary mismatch problem of an IR system. In this said problem, the query words mismatch with the document words. For example, when a user input a query word and the word does not exist in the vocabulary of the document then it may cause unreliable result. To avoid this problem we have developed a new web-based hybrid stemmer using brute force with enhanced suffix stripping algorithm union that can be adopt in the Odia information retrieval system. The new stemmer is both computationally faster as well as domain-independent.

II. RELATED WORKS

In the study of information retrieval, researchers find stemming plays an important role. Stemming is not a new concept. Stemming techniques had invented since 1968. The first stemming algorithm was designed by Julie Beth Lovins[1]. After that many researchers continued investigating various approaches to this area of study and proposed several algorithms to improve its performance.



Another stemmer in English was written by Martin Porter [2] in the year 1980. As compared to European languages as well as English,

a few researches have been discovered in Indian Language.

A Hindistemmer [3] was proposed by Rao, Durgeshet al. based on suffix stripping approach.

A Bengali Morphological analyzer [4] was developed by Dasgupta et al. based on suffix stripping approach. Stemming is the process by which the user inputs an inflected word to the trained model and the model produces the root/stem word according to its rule set. In this Paper we have developed a Stemmer based on Hybrid Approach.

III. LITERATURE SURVEY ON ODIA STEMMER

Reference	Key Findings
Sampa et al. [5]	Published a paper on Stemmer for Odia language. They used the suffix stripping approach to remove the inflectional suffixes. The limitation of this algorithm was it only predicts 88% accuracy.
Balbantray, R.C. et al. [6]	Published a paper FIRE 2012 Submission: MET Track Odia. They had used the affix removal algorithm. The system reads input text files from the folder. Firstly, it removes stop words from the input files against the stop word dictionary then matched the token with the root word dictionary. After that the input matched the suffix dictionary then removes the suffix and match with the root word. If the root word found then there is no further processing required.
Balabantaray, R.C. et al. [7]	Presented a paper on Odia Text Summarization.
Sethi, Dhabal Prasad [8]	Published a paper on Lightweight Stemmer for Odia Derivational Suffix. He used suffix stripping method to find the stems.

IV. ODIA DERIVATIONAL MORPHOLOGY

a) Odia Morphology:

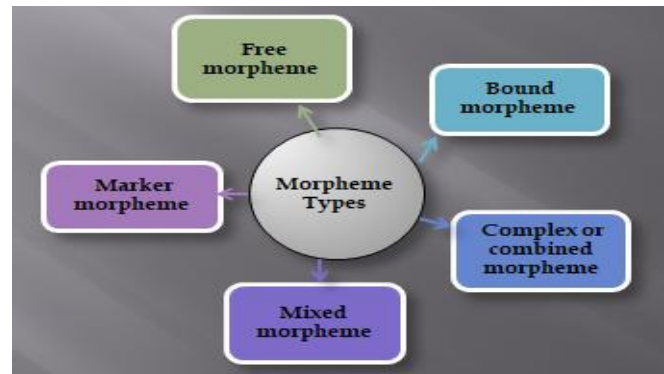
The formal variants of a morpheme are called allomorphs of that morpheme. The variant may be phonologically or morphologically conditioned. A morpheme may be a free or a bound form. Alternatively we can say that a word consists of one or more than one morpheme. From the point of view of its internal structure, a word may consist of (i) a root morpheme only (ii) a root and one or more non root morpheme or (iii) more than one root morpheme. The non-root morphemes are bound forms and are generally referred to as affixes. Roots enter into further morphological constructions and form a base while non-roots do not. [9]

Odia morphology deals with the analysis, identification and description of structure of morpheme. Morphology deals with the structure of words. The basic unit is the focus of study in morphology is morpheme.

Example: The word ବାଳକମାନେ the morphemes are ବାଳକ, ମାନେ. Morpheme is not always conveying a meaningful

word in Odia. Any morpheme in Odia should be a root word, prefix or suffix. Morphemes are divided into five categories shown in fig 1.

Fig. 1.Types of Morpheme



The morpheme which are independent called free morpheme. Those morphemes are standalone in nature. It does not need to add with other to create a word. Examples of free and bound morpheme discussed below.

ବାଳ ବାଚ ଖାଉଛି । ବାଳ ବାଚ(କୁ) ଖାଉଛି।

The morpheme ବାଚ is a stand-alone morpheme and morpheme (କୁ) is a suffix. Most of the morphemes are bound type in Odia language.

An affix is a morpheme that may come at the beginning (Termed as Prefix) or the end (Termed as Suffix) of a base morpheme. In Odia, prefixes are bound morphemes are affixes that come before a base morpheme. For example:

/ଉପକୁଳ/ = /ଉପ/ + /କୁଳ/

b) Odia Derivational Morphology:

Odia morphemes of different types such as nouns, verbs, affixes, etc. combine to form new words. The suffixes can be added with root word to form POS shown in the below Table- I.

Categories	Examples
Noun to Adjective	Noun word + Derivational suffix = Adjective category. ରୁପ+ଂଗୁଳି = ରୁପଙ୍ଗୁଳି
Adjective to Noun	Adjective word + derivational suffix = noun words ଆଧୁନିକ + ତା=ଆଧୁନିକତା
Adjective to Adjective	Adjective word + Suffix= Adjective word ସରୁ+ଆ=ସରୁଆ

Verb to Adjective	Verbal word + Derivational suffix = Adjective word ପାଳ+ଇଅ=ପାଳିଅ
Verb to Noun	Verbal word+Derivational suffix= Noun word. ବିଷ୍ଟର+ଅଣା=ବିଷ୍ଟରଣା ଖଳେ+ଆଳି=ଖଳୋଳି

In the suffix-stiffing approach of stemming, we remove suffixes or affixes attached to that stem. As Odia is being influenced by Sanskrit. So Odia language has strong inflectional system.

According to the Panini grammar rules, we are representing some rule in below example.

For example ‘କଲମଗୁଡ଼ିକ’ here stem କଲମ and suffix is ଗୁଡ଼ିକ .

2. The nominal suffix and verbal suffix of Odia language are described below. (Table – 1&2) [10].

Table- II: List of Nominal Suffixes in Odia

ବିଭକ୍ତି Inflection	ଏକ ବଚନ (Singular)	ବହୁ ବଚନ (Plural)	କାରକ (Case-Relationship)	ଉପପଦ (Noncase-Relationship)
ପ୍ରଥମା (1 st Inflection)	ଏକ,ଟକେ,ଟି,ଟା,ଟିଏ,ଟାଏଓଓ	ଏଗୁଡ଼ିଏ,ଗୁଡ଼ାକ,ଗୁଡ଼ିକ,ମାନ,ମାନକେ,	କରତା	
ଦ୍ୱିତୀୟା (2 nd Inflection)	କୁ,ଟକେ,ଟିକୁ,ଟାକୁ,ଝକୁ,କି,ଝକୁ, ଠାକୁଠିକି,ଠିକୁ,	ଝକୁଗୁଡ଼ାକୁ,ଗୁଡ଼ିକୁ,ମାନଝକୁ,	କରମ	
ତୃତୀୟା (3 rd Inflection)	ରଝେଝକଦ୍ୱାରା,ଦଲେ,ଦ୍ୱାରା,	ଝକଦ୍ୱାରା,ଦଲେ ମାନଝକ,ଦ୍ୱାରା ମାନଝକ, ମାନଝକରକେ	କରଣ	
ଚତୁର୍ଥୀ (4 th Inflection)	କୁ,ଝକ,ଲାଗି,ପାଇଁ,ଝକ,କି,ଝକୁ, ନିମନ୍ତକେ ଝକ,ନିମନ୍ତକେ ଝକ,ଲାଗି ଝକ, ସକାଶକେ ଝକ,ସକାଶକେ,ଯକେ-।ଗୁଁ	ଝକୁ,ଲାଗି ମାନଝକ,ଗୁଡ଼ାକୁ,ଗୁଡ଼ିକୁ,ମାନଝକୁ, ନିମନ୍,ମାନଝକନିମନ୍ତକେକେଯକେ-।ଗୁଁ ମାନଝକ, ସକାଶକେ	ସମ୍ପର୍କଦାନ	
ପଞ୍ଚମୀ (5 th Inflection)	ରୁଠଉଁ,ଠୁଁ,ଠାରୁ,	ମାନଝକଠାରୁ,ମାନଝକଠୁଁ, ଝକଠାରୁଠୁ ମାନଝକ,ଠଉଁ ମାନଝକ,	ଅପାଦାନ	
ଷଷ୍ଠୀ (6 th Inflection)	ର,ଠି,ଠାଇଁ,ଏ,ଝକର,ଝକ, ଠଲେଠି,ଠଉଁ,	ମାନଝକଝକରି,ମାନଝକରି,ମାନଝକର, ମାନଝକଝକରି,		ସମ୍ପର୍କଦ୍ୱୟ
ସପ୍ତମୀ (7 th Inflection)	ରକେ,ଠଲେ,ଠି,ଠାଇଁ,ଏ,ଠାରକେ, ଠଉଁଠି,	ମାନଝକରକେ,ଠାଇଁ ମାନଝକ,ମାନଝକଠାରକେ, ଠଲେଁ ମାନଝକ,ଠି ମାନଝକ	ଅଧିକରଣ	

Table- III: Odia Verbal Suffix

କାଳ (Tense)	ପୁରୁଷ (Person)	ଏକ ବଚନ (Singular Suffix)	ବହୁ ବଚନ (Plural Suffix)
ବର୍ତ୍ତମାନ କାଳ (Present Tense)	ପୁରଥମ ପୁରୁଷ	ଉଅଛି, ଉଛି, ଉଆନତି, ଏ	ଉଛୁ, ଇଅଛୁ, ଉଆନତୁ, ଉ
	ଦ୍ୱିତୀୟ ପୁରୁଷ	ଉଅଛୁ, ଉଛୁ, ଉଆନତୁ, ର	ଉଅଛ, ଉଛ, ଉଆନତ
	ତୃତୀୟ ପୁରୁଷ	ଉଅଛି, ଉଛି, ଉଅଛନ୍ତି, ଉଆନତା, ଉଆନତେ, ଉଆଆନତେ, ନତି, ଉଆନତା	ଉଛନ୍ତି, ଉଅଛନ୍ତି, ଉଆଆନତି, ନତି
ଅତୀତ କାଳ (Past Tense)	ପୁରଥମ ପୁରୁଷ	ଇଲି, ଇଛି, ଇଥିଲି, ଥିଲି, ଇଆନତି, ଇଅଛି	ଗଲୁ, ଉଥିଲୁ, ଉଥିଲୁ, ଇଆନତୁ
	ଦ୍ୱିତୀୟ ପୁରୁଷ	ଇଲୁ, ଇଛୁ, ଉଥିଲୁ, ଉଥିଲୁ, ଇଆନତୁ	ଇଲ, ଇଛ, ଇଅଛ, ଇଥିଲ, ଉଥିଲ, ଇଆନତ
	ତୃତୀୟ ପୁରୁଷ	ଇଲା, ଇଲେ, ଇଛି, ଇଛନ୍ତି, ଇଥିଲେ, ଇଥିଲଲ, ଉଥିଲେ, ଇଆଆନତା, ଇଆଆନତେ	ଇଲେ, ଇଛନ୍ତି, ଇଥିଲେ, ଉଥିଲେ, ଇଆଆନତେ
ବିବିଷ୍ଣୁତ କାଳ (Future Tense)	ପୁରଥମ ପୁରୁଷ	ଇବି, ଇଥିବି, ଥାନତି, ଉଥିବି	ଇବୁ, ଇଥିବୁ, ଥାଆନତୁ, ଉଥିବୁ, ଇବା
	ଦ୍ୱିତୀୟ ପୁରୁଷ	ଇବୁ, ଇଥିବୁ, ଥାଆନତୁ, ଉଥିବୁ	ଇବ, ଇଥିବ, ଥାଆନତୁ, ଉଥିବ
	ତୃତୀୟ ପୁରୁଷ	ଇବ, ଇବେ, ଇଥିବେ, ଇଥିବ, ଥାଆନତ, ଥାଆନତେ, ଉଥିବ, ଉଥିବେ	ଇବେ, ଇଥିବେ, ଥାଆନତେ, ଇଥିବେ

There are some prefixes which are attached to noun only in Odia language. There is 20 such prefixes exist in Odia. These are inherited from Sanskrit. These are as follows (Table -3):

Table- IV: List of Odia Prefixes

ପ୍ର	ପରା	ଅପ	ସମ୍	ନ
ଅଧି	ସୁ	ନିର୍	ଉତ୍	ପରି
ପ୍ରତି	ଅବ	ଅନ୍ନ	ଦୁର୍	ବି
ଅଭି	ଅତି	ଅପି	ଉପ	ଆ

V. PROPOSED METHODOLOGY FOR ODIA STEMMER

We have proposed a new web-based stemmer based on hybrid approach. [11][12] for Odia language. The new born stemmer is both computationally faster and domain independent. The algorithm of the proposed stemmer is described in the below flowchart. We developed a stemmer algorithm for Odia which removes derivational suffixes from derived word. The algorithm is a combination of brute force approach and enhanced suffix stripping approach. In this new enhanced model of suffix removal technique, system remove the suffixes recursively first 3 character, then 2 characters and last 1 character from the derived words.



VI. RESULT & DISCUSSION

We have evaluated the stemmer by taking different set of words i.e. 100 words, 200 words, 300 words and so on to calculate the time taken to extract the root words. We have not compared our Odia stemmer with any of the existing stemmer available for Odia language. Nowhere had we found the existing result to compare with the proposed stemmer.

Table 4: Time Taken to Extract Odia Root Words

Set No	No of Words	Time Taken (Sec.)
Set-1	100	5.07
Set-2	200	8.92
Set-3	300	13.27
Set-4	400	17.18
Set-5	500	20.29
Set-6	600	24.84
Set-7	700	29.11
Set-8	800	32.85
Set-9	900	37.66
Set-10	1000	40.48

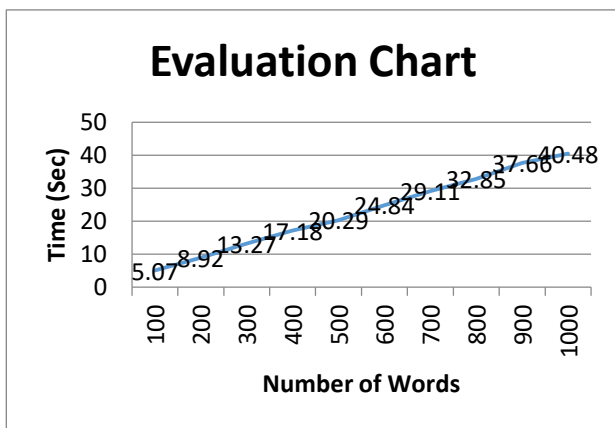


Fig. 3. Evaluation graph

VII. CONCLUSION AND FUTURE WORK

This stemmer can be used effectively to improve the query translation performance for information retrieval system [15]. Finally, we have integrated the new stemmer in the existing Dspace for Odia text retrieval system. The results are commendable and implicates that this can be used effectively in Odia Search Engine. It also handles the problem of under stemming and over stemming in some extent. In future we can integrate this stemmer for Odia text summarization purpose. As a future scope this system can be enhanced by including some more inflected word and corresponding to their stem mapping in to the database.

ACKNOWLEDGMENT

We are thankful to all the Language Researcher who directly and indirectly help to design and develop this web based Odia Stemmer.

REFERENCES

1. Lovins, J.B., "Development of a stemming algorithm", Mechanical Translation and Computational Linguistics, 1968, Vol. 11 Nos 1 and 2, pp. 22-31.
2. Porter, M.F. (1980), "An algorithm for suffix stripping", 1980, Program, Vol. 14 No.3, pp.130-137.

3. Ramanathan A, Durgesh D Rao, "A lightweight stemmer for Hindi", National Center for Software technology.
4. Dasgupta, S., V. Ng., "Unsupervised morphological parsing of Bengali", Lang Resources & Evaluation, Human Language Technology Research Institute, University of Texas at Dallas, Richardson, TX 75083, USA.
5. Chaupatnaik, S., Nanda, S.S., Mohanty, S., "A suffix stripping Algorithm for Odia Stemmer", Int. Journal of Computational Linguistics and Natural Language Processing, 2012.
6. Balabantaray, R.C., Sahoo, B., Swain, M., Sahoo, D.K., "IIIT-BH FIRE 2012 Submission: MET Track Odia".
7. Balabantaray, R.C., Sahoo, B., Swain, M., Sahoo, D.K., "Odia Text Summarization using Stemmer", Int. Journal of Applied Information Systems (IJ AIS), Foundation of Computer Science FCS, New York, USA, Volume 1- No.3, February 2012
8. Sethi, D.P., "Design of Lightweight Stemmer for Odia Derivational Suffixes", International Journal of Advanced Research in Computer and Communication Engineering Vol. 2, Issue 12, December 2013.
9. Odia grammar book of class 9th in BSE, ODISHA.
10. Chaupatnaik, S., Nanda, S.S., Mohanty, S., "A suffix stripping Algorithm for Odia Stemmer", Int. Journal of Computational Linguistics and Natural Language Processing, 2012.
11. Mohammad, A., Oqeili, S., and Rawan A. A., "Occurrences Algorithm for String Searching Based on Brute-force Algorithm", 2006 Jordan Journal of Computer Science Vol No.2, Issue No.1, pp 82-85.
12. Mishra U. et al., "MAULIK: An Effective Stemmer for Hindi Language", International Journal on Computer Science and Engineering (IJ CSE).
13. Erritali, M., "Information Retrieval: Textual Indexing Using an Oriented Object Database", Indonesian Journal of Electrical Engineering and Computer Science, Vol. 2, No. 1, April 2016, pp.205-214.
14. Sethi, D.P., "Morphological Analyzer for Sambalpur Odia Dialect Inflected Verbal Forms", International Journal of Advanced Research in Computer Science and Software Engineering, October, 2013.
15. Bajpai, P., Verma, P., Abbas, Syed Q., "Two Level Disambiguation Model for Query Translation", International Journal of Electrical and Computer Engineering (IJECE), Vol. 8, No. 5, October 2018, pp. 3923-3932.

AUTHORS PROFILE



Gouranga Charan Jena, is continuing his Ph.D. in CSE at School of Computer Engineering, KIIT DU, Bhubaneswar, Odisha, India. He has completed his ME(CSE) from Utkal University, Bhubaneswar, Odisha, India.

Area of Research: NLP, image processing, information retrieval, data mining, ML and Big data technologies.



Siddharth Swarup Rautaray, is working as an Assistant Professor in School of Computer Engineering, KIIT DU, Bhubaneswar, and Odisha, India. He has completed his Ph.D. in CSE from IIIT, Allahabad.

Research Interest: HCI, NLP, image processing and Big data technologies. He has published many international journals and conference papers on above research areas.