

Yorùbá Verb Sense Disambiguation using Semantic Similarity between Case Sentences of a Sense Inventory

Adegoke-Elijah, A.

Abstract The development of a word sense disambiguation component of a machine translation (MT) system is faced with many challenges. One of these is the incidence of contextual tonal variation in the Yorùbá language which makes the use of statistical based approach highly expensive for resolving lexical ambiguity in the language. This study examined the procedures underlining the resolution of lexical ambiguity in the context of Yorùbá-to-English MT system and developed a knowledge based approach which makes use of path-based similarity measurement between two instances of an ambiguous word to determine its right sense. This model achieved an accuracy of 96.1% for transitive verbs, 90.2% for intransitive verbs, with an overall accuracy of 94.6%, which is comparable with the high-performing supervised WSD, and a coverage of 69.3%. This study suggests a method that can be used to address ambiguity resolution in other low-resource languages.

Keywords: Sense, disambiguation, Yorùbá, similarity, ontology, sense inventory

I. Introduction

One of the distinguishing features between natural and formal languages is ambiguity [1-2]. The computational task of resolving lexical ambiguity that occurs in human languages within a given context is described as Word Sense Disambiguation (WSD) [3-5]. WSD is an important and challenging domain [6]. It is a computational task needed in many natural language tasks such as machine translation and sentiment analysis [7]. In the context of a machine translation, WSD can be described as the task of determining the right translation of a source language word, when the target language offers more than one possible translation in a bilingual dictionary [8-9]. Apart from English Language, WSD tool has been developed for many world languages, including Bengali [10] and Punjabi [11] etc. Despite the Yoruba language has over 30 million speakers residing within and outside Africa, there is dearth of studies in the development of ambiguity

resolution system for the language, and this has therefore limited studies in the development of an Yorùbá to English machine translation system, and vice-versa. Yoruba can be classified as a resource scarce language, primarily due to the difficulty in its orthography. It is a tonal language which adds diacritics to each syllables found in a word, to depict pitch pattern, and therefore aids the pronunciation of Yoruba words. The three (3) major tones used in the language are represented with grave (`), acute (/) and no symbol () for low, high and mid-tone respectively. The language also makes use of under dot to represent phonetic patterns in the vowels. As such the words *òkò* (farm) is different from *ókò* (husband), with the two words differentiated by under-dots in the vowel sounds. In the same way, *ìkòkò* (clay pot) and *ìkókó* (infant) are differentiated with the tone marks. Apart from these, there is also the incidence of interaction between tone and syntax in Yoruba sentences which makes low-tone monosyllabic verbs change to mid low when

Adegoke-Elijah, A.

(Department of Computer Science, Redeemer's University,
Ede, Nigeria)

Corresponding Author: adegoke-elijah@run.edu.ng

used in certain contexts. For example, the low tone verb *lù* (beat) changes to mid tone when used with a noun object as in *Adé lu ìlù náà* (Ade beats the drum), but retains its tone when used with a pronoun object as in *Adé lù ú* (Ade beats it). All these contribute to the difficulty in the availability of large amount of machine readable resources and sense-tagged corpora for the language. [12] did a study on the machine translation of English to Yorùbá texts. The aim of the work was to develop a system that could translate modified and non-modified simple English sentences. Though, the reported system accuracies were close to a human expert, the study did not address the translational ambiguity present in the language pair. [13] investigated the appropriate method needed for the disambiguation of verb senses in the Yorùbá language. The use of semantic role of the direct object of the ambiguous verb was found useful, and thus was used in a selectional preference method developed into hand-coded rules for the disambiguation task. The study showed that like in many languages, the use of hand-coded rule is laborious and often limits the flexibility and the scalability of WSD systems.

This study developed a verb disambiguation system which makes use of a knowledge based approach hinged on the semantic roles of the subjects and objects of the ambiguous verbs, and eliminates the need for hand-coded rules as witnessed in [13] and also circumvents the expensive and laborious task required in building sense-tagged corpora for the language.

Several methods have been adopted in resolving ambiguity in natural languages. Knowledge-based methods are based on algorithms that exploit the structural and lexico-semantic information found in external lexical resources; examples of such resources are machine readable dictionaries,

thesaurus, ontologies and WordNet, with Wordnet being one of the most widely use lexical resources in the field of natural language processing application [14-15]. The knowledge based approach has the advantage of not relying on the use of sense-annotated corpus which is labour intensive and not available for many resource scarce languages. Some of the approaches used in knowledge based methods are Lesk Algorithm [16], semantic similarity measurement [17], the use of selectional preferences [18] and heuristic based approaches. Unlike other approaches, knowledge based techniques can disambiguate more than one target word at the same time, and disambiguate them jointly; have been shown to achieve competitive results [6], [19-20].

The supervised approach to WSD makes use of sense annotated corpora to train classifier to determine the right sense of a word using machine learning algorithm. Some of the techniques used in supervised WSD are decision list, decision tree, Naive Bayes, the use of neural network, support vector machine [3], [21]. In general, supervised WSD demonstrate better performance than other techniques, its demerit is its reliance on sense-annotated dataset, which is highly expensive to build and not practicable to handle all the words in a language due to the need to build separate word expert for each word to be considered.

Unsupervised WSD does not rely on sense annotated dataset, but aim to determine the right sense of a word using clustering based measurement of contextual similarity, hence addresses the challenge of knowledge acquisition found in supervised WSD approach. Techniques for unsupervised WSD include context clustering, word clustering and the use of concurrence graph [22]; [23-24].

The advent of pre-trained models has demonstrated promising results in many natural language processing tasks, including word sense disambiguation. This is based on the ability of the model to capture rich contextual information and also offer high-quality word representation, which can be used for WSD in many ways. One of these is for fine-tuning of labelled WSD dataset. The pre-trained language model can also be utilized to capture the contextual meaning of a word through the generation of contextualized word embedding. The model can be used to generate sense embedding for ambiguous words. Pre-trained models have been adopted for several WSD tasks including [25-27].

The task of Verb Sense Disambiguation (VSD) involves assigning the right sense to an ambiguous verb automatically [9], [28]. It is thus a sub-problem in word sense disambiguation. All verb sense disambiguation methods depend on the features of the words found in the context of the ambiguous verb; however, semantic feature plays a very important role in the disambiguation of verbs [29]. Verb sense disambiguation systems often depend on the distinctions in the semantic of the target verb's arguments. In the same vein, [30-31] stated that predicate-argument information and selectional restrictions are hypothesized to be particularly useful for disambiguating verb senses.

This study presents a knowledge base approach that made use of two external resources for the disambiguation of Yoruba verbs. The external resources are hand-crafted ontology that depicts *is-a* relationships among Yoruba nouns, and a lexical database of Yoruba verbs. The relationships between nouns shown in the ontology provides the semantic feature used for the verb disambiguation, while the lexical database provides case sentences for the

measurement of semantic similarity between concepts. The knowledge-based approach was adopted because of its ability to handle multiple ambiguous words jointly and its usability for the disambiguation of verbs in resource-scarce languages, like the Yoruba language.

II. Materials and Methods

A. Task Description

Given a Yorùbá sentence Y comprising a sequence of words W_i , with an ambiguous word w_0 . Given that the ambiguous word is w_0 , and that i is a counter denoting the relative distance of the contextual words to the ambiguous word.

$w = \{w_i | i < 0\}$ are words that at the left hand side of the ambiguous word;

$w = \{w_i | i > 0\}$ are words that are at the right hand side of the ambiguous word.

Given also a set $T = \{t_1, t_2, \dots, t_n\}$ contains the possible translations of w_0 in English language, where n is the number of possible translations of w_0 as described in the sense inventory. The task of the WSD is to select the right translation of w_0 , from the set T using the features of w_i .

This is represented by the mapping function W to T , $W \times P \rightarrow T$. That is, the proposed system chooses the right translation for w_0 from set T using the elements of P , where P is a set containing the semantic features of the nouns and pronouns found in the set W .

Using the sentence *Mo fún Táyé ní owó*, as an illustration,

$Y = \{Mo, fún, Táyé, ní, owó, \}$

Where w_0 is the ambiguous word *fún*,
 $w_{-1} = Mo$, $w_1 = Táyé$, $w_2 = ní$, $w_3 = owó$

$T = \{give, for, tight\}$

P

$= \{feature\ of\ Mo, feature\ of\ Táyé, feature\ of\ owó\}$

The task of the proposed system is to choose the appropriate translation for *fún* from the set *T* using the elements of set *P*.

B. Methods

To achieve the task described above, this study made use of distributional hypothesis [32-33] and uses the nouns (arguments) found in the context of the ambiguous verb to determine the correct sense of the verb. This is based on semantic similarity between concepts. [34] defined semantic similarity between concepts based on Quillian spreading activation theory [35]. One of assumptions of spreading activation theory is that semantic network (ontology) is organized along the line of similarity. The conceptual distance, which is used to quantify the similarity between concepts, could be measured as the geometric distance between two points representing the concepts. The conceptual distance between two concepts is a decreasing function of the similarity, which means the more similar two concepts are, the smaller the conceptual distance between them. [34] computes the semantic distance between the concepts by counting the number of edges between them in the ontology.

Let C_1 and C_2 be the two concepts in an is-a semantic network. The conceptual distance between C_1 and C_2 is given by:

$$Distance(C_1, C_2) = \text{Minimum number of edges separating } C_1 \text{ and } C_2 \quad (1)$$

$$Pathlen(C_1, C_2) = 1 + Distance(C_1, C_2) \quad (2)$$

$$Simpath(C_1, C_2) = \frac{1}{pathlen(C_1, C_2)} \quad (3)$$

where *pathlen()* is a function returning the path length (number of edges) between C_1 and C_2

and *simpath()* is the similarity measure between the two concepts.

The steps taken to achieve the stated objectives are described in the following sections.

C. Data Collection

This involves the collection of data needed for the analysis of the proposed model. Two types of data were collected for this study. These are verb data and noun data. The process of collecting and preparing each of these data is described below:

i. Verb data

The verb data used in this study is the one reported in [36]. It is made up of 93 verbs tagged with their various translations in English language. It also contains sample usage sentence for each of the senses in Yoruba language. A screen shot of the entry for the ambiguous verb *bó* which has four possible translations in English language is shown in Figure 1.

```
<entry entry_id="06" word= "bó"senses="4">
  <sense sense_id ="1"
    translation="pull-off"
    case= "Tolú bó asọ rẹ"></sense>

  <sense sense_id ="2"
    translation="drop"
    case= "Ó bó sí ilẹ"> </sense>

  <sense sense_id ="3"
    translation="feed"
    case= "Adé bó ọmọ náà"> </sense>

  <sense sense_id ="4"
    translation="enter"
    case= "Ó bó sí yàrá"> </sense>

</entry>
```

Fig. 1: Sample Verb Lexical Database

ii. Noun Data

A total number of three hundred and two (302) Yorùbá nouns were collected from the home domain and organized into sub-class hierarchy called an ontology as shown in Figure 2. Using graph theory, the ontology can be described as a

rooted graph with twenty-four (24) vertices and twenty-three (23) edges. The twenty-four vertices are the semantic categories of the collected nouns arranged in subclass hierarchy. The ontology consists of seven (7) internal vertices Top, Concrete, Non-living thing, Living thing, Location, Solid and Human, with the vertex top representing the superclass. It is also made up of seventeen (17) leaf nodes which are Abstract, Wears, Properties, Woods, Food, Tools, Liquid, Gas, Plant, Body parts, Profession, Names, Kingship, Pronoun, Animal, common and proper nouns.

Each of the leaf nodes can be described as a finite set that contains noun that are of its kind. For example, the *Food* node is a finite set containing items which are names of food. The main aim of developing this ontology is the measure the semantic similarity between two nouns(concepts) found on the ontology. The possible number of edges between two concepts in the ontology is between one (1) and seven (7). Using Equation (2), the maximum measurement of semantic similarity between concepts in the ontology is 0.5, while the minimum semantic similarity measurement is 0.125

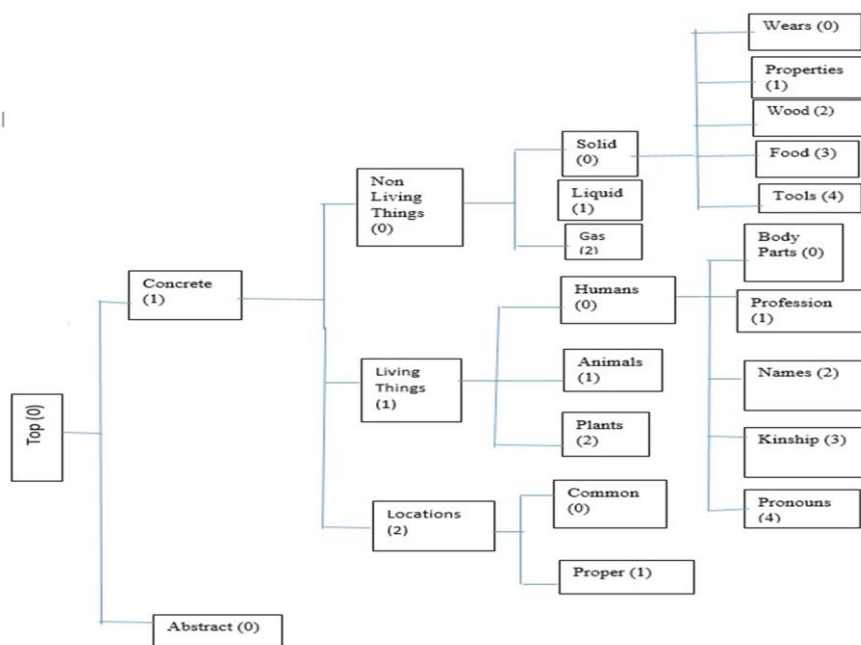


Fig. 1: Hierarchical classification of Yoruba Nouns

D. Model Formulation

The proposed model is based on path-based semantic similarity measurement between two sentences. As described earlier, the sense inventory is made up of the ambiguous verbs, their possible translations in English language and usage example sentences. To choose the

right sense of an ambiguous verb in a Yorùbá test sentence, the disambiguation algorithm measures the semantic similarity between a given test sentence and each of the case sentences corresponding to the verb as described in the sense inventory. This is achieved by summing up the similarity of all the corresponding arguments in the sentences. The case sentence with has the

greatest similarity with the test sentence is chosen as the right sense of the verb, and the corresponding English language translation of the Yorùbá verb is displayed as the correct translation of the ambiguous verb. The formal description of the disambiguation process is given below.

Given Yorùbá test sentence T , containing ambiguous verb W_0 with arguments $t_1, t_2, t_3, \dots, t_m$. Given also case sentences C_i which are example sentences of w_0 derived from the sense inventory, where $i = 1$ to n , and n is the number of possible translations of W_0 . C_i contains argument $C_{i,1}, C_{i,2}, \dots, C_{i,m}$. Where m is the total number of arguments found in the context of the ambiguous verb. The equation for calculating the similarity between T and each of the case sentences C_i is :

$$Simpath(T, C_i) = \sum_{j=1}^m \frac{1}{P(t_j, C_{ij})} \quad (4)$$

Where S and P are the *Simpath* and *Pathlength* functions respectively.

$$S_i = \max_{1 \leq i \leq n} S(T, C_i) \quad (5)$$

The sense S_i which has the greatest similarity with the test sentence T is therefore chosen as the right sense of the ambiguous verb.

To illustrate the formal method described above, the task of disambiguating the Yorùbá verb $bò$ which has four senses as specified in the sense inventory. The possible translations are *drop, feed, enter and pull-off*. Given a test sentence $\dot{O}jò bọ̀ Sòkòtò rẹ̀$, the proposed method chooses the right translation for the verb from the list of the possible translations by measuring the semantic similarity between the test instance and each of the case sentences $\dot{O} bọ̀ sí ilẹ̀$, $Tolú bọ̀ asọ̀ rẹ̀$, $Adé bọ̀ ọmọ náà$ and $\dot{O} bọ̀ sí yàrá$ corresponding to the

verb. The features extracted from each of the sentences are the noun arguments found in the context of the ambiguous verb in each of the instances. The arguments extracted from the test sentence are $[\dot{O}jò, Sòkòtò, rẹ̀]$. The same features are extracted from each of the case sentences. The features extracted case sentences 1, 2, 3 and 4 are $[\dot{O}, ilẹ̀]$, $[Tolú, asọ̀, rẹ̀]$, $[Adé, ọmọ]$ and $[\dot{O}, yàrá]$ respectively.

The next step is to calculate the semantic similarity between each of the features extracted from the test sentence (test features), and the features extracted from the case sentences (case features). That is, the proposed method calculates the semantic similarity between $[\dot{O}jò, Sòkòtò, rẹ̀]$ and $[\dot{O}, ilẹ̀]$ for sense one, the similarity between $[\dot{O}jò, Sòkòtò, rẹ̀]$ and $[Tolú, asọ̀, rẹ̀]$ for sense two, the similarity between $[\dot{O}jò, Sòkòtò, rẹ̀]$ and $[Adé, ọmọ]$ for sense three, and finally calculate the similarity between $[\dot{O}jò, Sòkòtò, rẹ̀]$ and $[\dot{O}, yàrá]$ for sense four. The semantic similarity between test sentence and each of the usage sentences are 0.58, 2, 0.125 and 1.142 respectively, and this therefore shows that usage sentence 2 has the greatest semantic similarity with the test sentence, and therefore its corresponding translation, *pull-off* is there for chosen for the verb $bọ̀$ in this context.

The model formulated above was reduced into an algorithm which accepts a Yoruba sentence, the developed sense inventory and the ontology as inputs, and outputs the translated verb as shown in Figure 3.

E. System Implementation

The algorithm was thereafter implemented using Python 3.5 programming Language. The language was chosen because it contains libraries for processing human language. The software

```

Data: Yoruba sentence, lexical database, ontology
Result: Translation verb

i, j = counter;
k, n = integer;
Y= Input Yoruba sentence;
Tokens= Tokenize(Y);
POS= POSTag(Tokens);
[c1, c2,...,ck]= Extractnoun(Y)
w0= verb;
if w0 in lexical database then
    n= no of senses;
    for j= 1 to n do
        Sj= case sentence;
        tokens= tokenize(Sj);
        POS= POSTag(Sj);
        [x1,x2,..., xk]= Extractnoun(Sj);
        similarity(Sj)= 0;
        for i= 1 to k do
            similarity(Sj) = similarity(Sj) + similarity(Ci,Xi);
        end
    end
    sense= max(Similarity(Sj));
    translate(w0);
else
    Print("Verb not in the database");
end

```

Fig 3: The Algorithm for the Disambiguation Process

contains some functions which are the building blocks of the developed software. *ExtractNoun()* outputs the nouns found in the context of an ambiguous verb. *SearchforAmbiguousVerb()* identifies the ambiguous verb in the Yoruba sentence by simply checking for its existence in the sense inventory developed for this study. *SearchforPossibleMeanings()* scans through the sense inventory and outputs the possible meanings of an ambiguous verb. *GetWordsDistance()* accepts two nouns and computes the number of nodes between them using the developed ontology. *SemanticSimilarity()* finds the inverse of the

distance between two nodes to calculate the semantic similarity. *FindMeaning()* finally outputs the correct meaning of the ambiguous verb by selecting the sense with the greatest similarity with the test Yoruba sentence.

The software developed is called Yorùbá Lexical Disambiguator (YoLeD). The implemented software takes in a Yorùbá text was containing an ambiguous verb as input, shows the possible meanings of the verb, and outputs the correct translation based on the context. A screenshot of the developed software is shown in Figure 4.

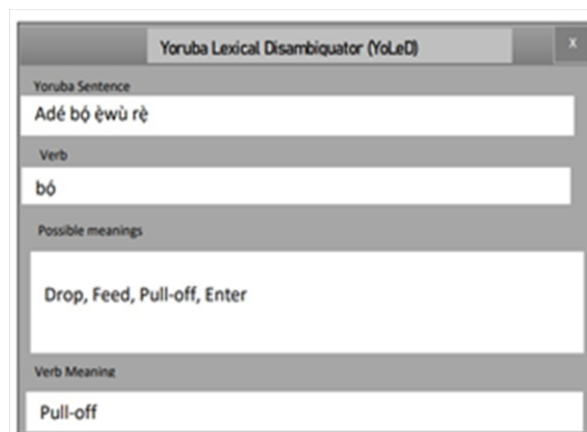


Fig 4: A Screenshot of the Implemented System

F. System Evaluation

The implemented system was evaluated using accuracy and coverage performance metrics. Details of the evaluation process, using each of these metrics, are discussed in the following subsections.

i. Accuracy

This is the metric used to measure the accuracy of the system in translating ambiguous words correctly in a particular context. Given a Yorùbá test sentence containing an ambiguous verb, the sentence is entered into the implemented system to perform the task of translating the verb in the sentence to English language.

Table 1: System Evaluation

	No of senses	Number of correct instances	Accuracy
No of senses (Transitive verbs)	180	173	96.1%
No of senses (Intransitive verbs)	61	55	90.2%
	241	228	94.6%

Given the total number of possible translations of the verb, the evaluation task checks if the result of the verb translation is correct in the context. Given a number of test sentences, the instances in which the verb translations are accurate are marked correct, while the instances in which the results of the verb translations are inaccurate are marked incorrect.

$$Accuracy = \frac{No\ of\ correct\ senses}{Total\ number\ of\ possible\ senses} \times 100 \quad (6)$$

To derive the total number of senses, the number of possible translations of all the verbs

considered in this study are summed together. Out of the total number of senses, the number of correct senses is determined by counting how many of the verb translations carried out by the system are marked correct. The summary of the results of the evaluation for the system is shown in Table 3.

ii. Coverage

In this study, coverage is defined as the percentage of the ambiguous monosyllabic verbs covered in the implemented system, in relative to the total number of ambiguous verbs in the Yorùbá language. That is:

$$Coverage = \frac{No\ of\ ambiguous\ verbs\ considered}{Total\ number\ of\ ambiguous\ verbs} \times 100 \quad (7)$$

According to Adegoke-Elijah (2018) which presented a method of estimating the total number of ambiguous verb in the language by using the grammar of Yoruba verbs, the total number of ambiguous monosyllabic verbs in Yorùbá language is one hundred and forty-one (141). Therefore,

$$Coverage = \frac{93}{141} \times 100 = 65.96\%$$

III. Results and Discussion

The developed system was able to disambiguate 96.1% of the transitive verb correctly. Such verbs have at least one argument that belongs to different semantic categories (classes) for each sense of the verb. For example, if we consider the ambiguous verb *dín*, with stores case examples, *Mo dín eja* for fry sense of the verb, and *Solá dín owó* for reduce sense of the verb. The two senses have at least one argument that belongs to different semantic classes i.e. *owó* belongs to property semantic category and *eja* belongs to food semantic category. This is in

agreement with [9] who hypothesized the usefulness of predicate-argument information and selectional restrictions for disambiguating verb senses. For intransitive verbs, the system achieves a lower accuracy of 90.2% for verbs in which the pathlength between the features in the case sentences and the test sentence is minimal; however, the accuracy of the system reduces as the pathlength between the feature increases. For example, in the case of the intransitive verb *bò* with the case sentence *Ọmọ náà ní bọ* assigned to the *come* sense of the verb in the sense inventory, the system was able to correctly translate the verb in the test sentence, *Bàbá ní bọ*, but incorrectly translated it in the sentence *Tolú ní bọ*. This is because the pathlength between *Tolú* and *Ọmọ* is 4, with the semantic similarity measurement of 0.25; whereas the path length between *Bàbá* and *Ọmọ* is 2, with semantic similarity measurement of 0.5.

The accuracy of the system could be improved if the number of case sentences used for each sense of the ambiguous verbs are more than one. This will give room for some verbs which can take more than one semantic category. For example, the ambiguous verb *yọ* with the remove sense was incorrectly translated in the test sentence *Adé yọ abọ́* to *appear* because only one case sentence *Adé yọ ẹxẹ* was used in the sense inventory. The system therefore expected an argument related to a living thing as the object of the verb, even though the argument *abọ́* which is a non-living thing can also be used in the *remove* sense of the verb.

IV. Conclusion

This study has presented an approach for addressing translational ambiguity in a Yoruba to English machine translation system. The method makes use of a knowledge-based approach that depends on the case (example) sentences

provided for each sense of an ambiguous verb, and also an ontology which provides the semantic features for the nouns found in the context of the ambiguous verb. This study concludes that the method can be adopted for other resource scarce language without sense tagged corpora.

References

- [1] Yadav, A., Patel, A., Shah, M."A Comprehensive Review on Resolving Ambiguities in Natural Language Processing." *AI Open*, Volume 2, 2021 pp. 85-92.
- [2] Sennet A. "Ambiguity. In *The Stanford Encyclopaedia of Philosophy*". CSLI Publications, 2015, <http://plato.stanford.edu/archives/spr2015/entries/ambiguity/>
- [3] Mishra, B. K., & Jain, S. "An Innovative Method for Hindi Word Sense Disambiguation" *SN Computer Science*, Volume 4, Number 704 2023.
- [4] Navigli. R. "Word Sense Disambiguation: a survey" *ACM Computing Surveys*, Volume 2009, pp. 10-64
- [5] Pal A. R., Diganga S. "Word Sense Disambiguation: A Survey" *International Journal of Control Theory and Modelling* Volume 5, 2015, pp. 1-15
- [6] Kokane, C. D., Babar, S. D., Mahalle, P. N., Patil, S. P. "Word Sense Disambiguation: Adaptive Word Embedding with Adaptive-Lexical Resource" *The International Conference on Data Analytics and Insights*, Singapore.2023, pp. 421-429.
- [7] Zhang, X., Mao, R., He, k., Cambria, E "Neuro-symbolic SENTIMENT Analysis with Dynamic Word Sense Disambiguation" *Association for*

- Computational Linguistics: EMNLP 2023*, Singapore, 2023, pp. 8772-8783.
- [8] Barman, A. K., Sarmah, J., Basumatary, S., Nag, A. "Word Sense Disambiguation applied to Assamese-Hindi Bilingual Statistical Machine Translation. *Engineering, Technology & Applied Science Research*, Volume 14, 2024, pp. 12581-12586.
- [9] Oliveira F, Wong F, Li Y, and Zheng J "Unsupervised Word Sense Disambiguation and Rules Extraction using non-aligned Bilingual Corpus", *The International Conference on Natural Language Processing and Knowledge Engineering*. 2005, pp. 30-35.
- [9] Das D., Khan, A., Shaikh, S. H., Pal, R. K. "A Dataset for Evaluating Bengali Word Sense Disambiguation Techniques" *Journal of Ambient Intelligence and Humanized Computing*, Volume 14, 2023, pp. 4057-4086.
- [10] Singh, V. P., Kumar, P. "Word Sense disambiguation for Punjabi Language using Deep Learning Techniques" *Neural Computing and Applications*, Volume 32, 2020, pp. 2963-2973.
- [11] Elúdióra S. and Qdéjobi O. "Development of an English to Yorùbá Machine Translator" *International Journal of Modern Education and Computer Science*, Volume 8, 2016, pp. 8-15.
- [12] Adégòke-Elijah A., Qdéjobi O. and Saláwú A. "Lexical Ambiguity Resolution for Standard Yorùbá Verbs" *American Journal of Engineering Research (AJER)*, Volume 7, 2018, pp. 170-176.
- [13] Kazeminejad, G. "Computational Lexical Resources for Explainable Natural Language Understanding" PhD diss., University of Colorado at Boulder, 2023.
- [14] Phye S. L. "Development of Myanmar-English Bilingual WordNet like Lexicon". *International Journal of Information Technology and Computer Science*, Volume 10, 2014, pp. 28-35.
- [15] Kumari, L., & Kumar, S. "Optimizing Word Sense Disambiguation for Hindi language using extended Lesk and Conceptual Density" *8th International Conference on Computing in Engineering and Technology (ICCET 2023)*, 2023, pp. 216 – 219.
- [16] Jha, P., Agarwal, S., Abbas, A. "A Novel Unsupervised Graph-Based Algorithm for Hindi Word Sense Disambiguation" *SN Computer Science*. Volume 4, 2023, pp. 675-684.
- [17] Ye B. and Baldwin T. "Verb Sense Disambiguation using Selectional Restriction extracted with a State-of-the-art Semantic Labeller", *Australasian Language Technology Workshop*, Sydney, Australia, 2006, pp 137-146.
- [18] Moro A., Raganato A., and Navigli R. "Entity Linking meets Word Sense Disambiguation: a Unified Approach" *Transactions of the Association for Computational Linguistics (TACL)*, 2014, pp. 231-244.
- [19] Agirre, E., Lacalle O. "Publicly Available Topic Signatures for all WordNet Nominal Senses." In *LREC*. 2004.
- [20] Alian, M., Awajan, A.:Arabic "Word Sense Disambiguation using Sense Inventories". *International Journal of Information Technology*, Volume 15, 2023, pp. 735-744.
- [21] Hou, B., Qi, F., Zang, Y., Zhang, X., Liu, Z., Sun, M. "Try to substitute: An unsupervised Chinese Word Sense

- Disambiguation Method based on HowNet” *Proceedings of the 28th International Conference on Computational Linguistics*, 2020, pp. 1752-1757.
- [22] Rahman, N., Borah, B. “An unsupervised method for word sense disambiguation” *Journal of King Saud University-Computer and Information Sciences*, Volume 34, issue 9, 2022 pp. 6643-6651.
- [23] Rahmani, S., Fakhrahmad, S. M., Sadreddini, M. H. “Co-occurrence graph-based Context Adaptation: a new Unsupervised Approach to Word Sense Disambiguation” *Digital Scholarship in the Humanities*, Volume 36, 2021, pp. 449-471.
- [24] Vandenbussche, P., Scerri T., Daniel R. “Word Sense Disambiguation with Transformer Models” *Proceedings of the 6th Workshop on Semantic Deep Learning (SemDeep-6)*, 2021, pp. 7-12
- [25] Duarte, J M., Samuel S., Evangelos M., Lilian B. “Deep Analysis of Word Sense Disambiguation via Semi-supervised Learning and Neural Word Representations” *Information Sciences*, 570, 2021, pp. 278-297.
- [26] Pasini, T., Alessandro R., Navigli R, “XL-WSD An extra-large and Cross-lingual Evaluation Framework for Word Sense Disambiguation” *In Proceedings of the AAAI Conference on Artificial Intelligence*, 2021, pp. 13648-13656.
- [27] Dutta, A., Samir K B. “Verb Sense Disambiguation by Measuring Semantic Relatedness between Verb and surrounding Terms of Context” *International Journal of Advanced Computer Science and Applications*, Volume 12, 2021.
- [28] Gung, J., Martha P. “Predicate Representations and Polysemy in Verbnet Semantic Parsing” *Proceedings of the 14th International Conference on Computational Semantics (IWCS)*, 2021, pp. 51-62.
- [29] Pal, A. R., Diganta S., Sudip K. N., Niladri S. D. “In search of a Suitable Method for Disambiguation of Word Senses in Bengali.” *International Journal of Speech Technology*, Volume 24, 2021, pp. 439-454.
- [30] Dang H. T, Chia C., Palmer M, Chiou F. “Simple Features for Chinese Word Sense Disambiguation” *Paper presented at the 19th International Conference on Computational Linguistics*, Volume 1, 2002, pp. 1-7.
- [31] Sahlgren M. “The Distributional Hypothesis”, *Italian Journal of Linguistics*, Volume 20 Issue 1, 2008, pp. 33-54.
- [32] Firth J. R. “A synopsis of linguistic theory: 1930-1955” *Studies in Linguistic Analysis*, 1957, pp. 1-32.
- [33] Rada R, Mili H, Bicknell E, and Blettner M. “Development and Application of a Metric on Semantic Nets” *IEEE Transactions on Systems, Man, and Cybernetics*, Volume 19, 1989, pp.17-30.
- [34] Quillian M. “Semantic Memory”. *Semantic Information Processing*. MIT Press, Cambridge. 1968.
- [35] Adegoke-Elijah, A, Jimoh K and Alabi A., “Development of an XML-Encoded Machine Readable Dictionary for Yoruba Word Sense Disambiguation” *Uniosun Journal of Engineering and Environmental Sciences*, Volume 5, 2023, DOI: 10.26108/ujees/3202.50.0210