

# Developing a Multilingual Stemmer for the Requirement of Text Categorization and Information Retrieval

Said Gadri<sup>1</sup> and Erich Neuhold<sup>2</sup>

<sup>1</sup>Department of Computer Science, Faculty of Mathematics and Informatics  
University Mohamed Boudiaf of M'sila, M'sila, Algeria

<sup>2</sup>Department of Computer Science, Faculty of Computer Science  
University of Vienna, Vienna, Austria  
Said.kadri@univ-msila.dz, erich.neuhold@univie.ac.at

**Abstract:** Information retrieval IR is the process of finding information (generally documents) that matches the needs of the user. One way to improve the search effectiveness, as well as the quality of text categorization is to build an effective stemmer that helps to match users' queries with relevant documents in IR and reduce the space of textual representation in TC. This has been always an interesting research topic in IR and TC. We can define stemming as the process of reducing inflected and derived words to their reduced forms (stems or roots). Many stemmers have been developed for different languages, but there is always many weaknesses and problems. In the present work, we have developed a multilingual stemming approach, based on the extraction of the word root and that exploits the technique of n-grams of characters. Our experiments have been done on three languages which are: Arabic, English, and French.

**Keywords:** Information retrieval, Machine learning, Natural language processing, Root extraction, Stemming

## 1. Introduction

Text categorization TC is an interesting task in the field of text mining. It is useful in several applications such as spam filtering, assignation of web pages to predefined groups (e.g. yahoo groups), etc. it becomes more and more important as the number of texts in electronic format grows every day. This task is done by assigning a set of texts to another set of predefined categories (classes). To fulfill TC requirements, there exist two approaches: the manual approach (human experts) which is focused on the manual development of classification rules, machine learning approach which is based on the use of algorithms known in the field of machine learning such as K-NN, NB, Decision trees, NB, SVM, ANN, etc. During TC process, the document must pass through many steps including: pre-processing step which consists of removing irrelevant words such as punctuation and stopwords, representation step which consists of representing each document with a vector of terms (words, phrases, n-grams, etc), and finally the calculation step in which we calculate term frequencies TF, and inverse document frequencies TF-IDF. Unfortunately, one critical problem can be raised during the representation step which is the large size of vectors used to represent these documents, this case occurs when we use big corpora of texts. To overcome this problem, a set of statistical methods can be used to select the most relevant terms and use them in the entry of learning algorithms [1-3]. Another technique that gives good results for the categorization of Arabic texts is the selection of relevant terms using stemming technique.

In the field of Information Retrieval IR, the stemming is also used to conflate a word to its different forms to avoid inconsistency between the words occurring in the documents and the query asked by the user. For example, if the user is looking for a document on "How to draw" and he submits a query on "drawing", he may not obtain all results he is looking for. But, if his query is stemmed, so that "drawing" becomes "draw", then retrieval will be more successful.

Received: October 14<sup>th</sup>, 2021. Accepted: June 1<sup>th</sup>, 2022

DOI: 10.15676/ijeei.2022.14.2.3

Thus, if we want to give a simple definition of stemming, we can say that stemming is a linguistic technique that permits to replace a large number of terms (words) occurring in a collection of documents and semantically close by their reduced forms (roots or stems), the main objective is to reduce the size of the vector of terms on one hand, and to improve the accuracy of categorization in TC and the effectiveness of search in IR on the other hand.

Several stemmers have been built for different languages such as English, French, Spanish, Italian, and Arabic. Each one has its weaknesses as well as its disadvantages. We note also, that most of these stemmers are language-dependent [4]. So, it seems useful to design and develop a new stemming algorithm that is totally language-independent. For this purpose, we developed in our study a multilingual stemming approach, that relies on the search of the word root. This approach has been validated on three different languages namely: English, Arabic, and French, and gave promising results. We also note that our proposed algorithm can work well for the lemmatization process, but we have focused here only on the stemming process which is our first concern. The present paper is organized as follows: Section 2 presents some concepts related to the stemming process. Section 3 is a detailed overview of the related works. In section 4, we describe our proposed stemming approach. The fifth section illustrates the experimental part we have done to validate our approach, as well as the obtained results. section 6 presents the major advantages of our new stemming approach. In section 7, we discuss the results obtained in the previous section. Section 8 displays a comparison established between our proposed multilingual stemmer and the most popular ones in the field. In the last section, we summarized the realized work and suggested some perspectives for future research.

## **2. Basic Concepts**

### *A. Stemming and Lemmatization*

Stemming and lemmatization have almost the same function. Both of them replace a word with its basic form; a 'root' or 'stem' for stemming and 'lemma' for lemmatization. Therefore, there exists a main difference between both concepts. In stemming the (root/stem) is obtained after applying a set of rules (removal of affixes) without requiring a grammatical analysis of the text. In contrary, the lemmatization process is consisting in obtaining the 'lemma' of a word which requires replacing conjugated verbs with their infinitive forms and nouns with their singular forms after applying a grammatical analysis on the text (e.g., POS tagging). So, it is a more complicated operation than stemming especially when implementing it on the machine.

### *B. Over Stemming and Under Stemming Errors*

In stemming, two errors may be appeared: over stemming and under stemming. We talk about over-stemming if two words having different stems are stemmed to the same root/stem. similarly, we talk about under-stemming if two words that should be stemmed to one root/stem are not. Both errors influence negatively the performance of the stemmer. We note also that, when a word is under-stemmed, it becomes enough difficult to determine if two different words are related according to their morphology or not. In the case of over stemming, the major problem is that it will be possible that two no related words share a common part of their morphological root are incorrectly detected as related because their stems are identical. Many researchers have proved that light-stemming decreases the over-stemming errors but increases the under-stemming errors. In contrast, heavy stemming reduces the under-stemming errors but increases the over-stemming errors [5, 6].

### *C. Conflation Between Word Variants*

The term conflation is generally used to denote the fact of matching morphological term variants. We distinguish two kinds of conflation methods; manual conflation based on some regular expressions and automatic conflation based essentially on some programs called stemmers. Figure 1 explains briefly the different methods of conflation.

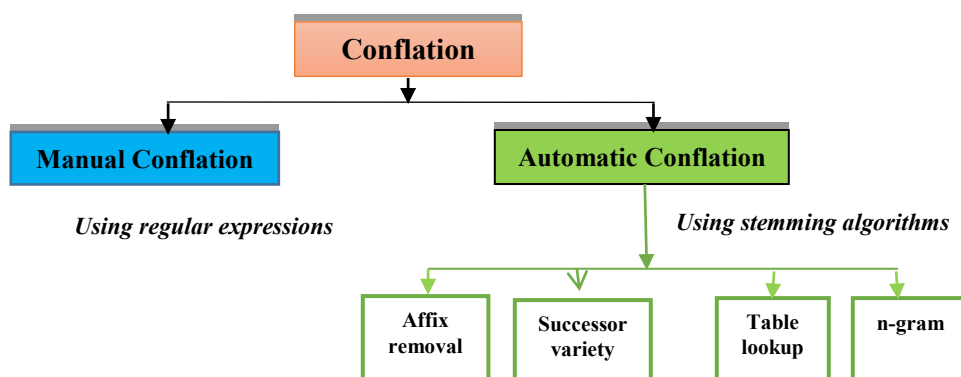


Figure 1. Conflation methods

*Affix removal methods:* This class of methods removes affixes (suffixes, infixes, prefixes) from the words subject of conversion into a common form called a stem. Most stemmers are based on the affix removal methods. They are also based on two main parameters: iterations and the longest match possible. We define an iterative stemming algorithm as a recursive process that removes strings in each order-class, starting from the end of the word toward its beginning. This approach of stemming does not allow more than one match within a single order class [7]. The meaning of longest-match is that within any given class of endings if more than one ending provides a match, the longest one should be removed.

*Successor variety method:* This class of methods generally uses the frequencies of letter sequences in a given body text as the basis of the stemming algorithm. In simple terms, the successor variety of a string can be defined as the number of different characters that follow it in words in somebody of text [8]. For example, we take the following body of text consisting of many words: (call, cake, cease, chair, cash, cost, cheap, core, check, cubic). To find the successor varieties for the word "carpenter" for example, we must use the following process: The first letter of "carpenter" is "c". "c" is followed by five characters in the text body: "a," "e," "h," "o," and "u." Thus, the successor variety of "c" is five (05). The next successor variety for "carpenter" would be three (03), since "ca" is followed by "l," "k," "s." in the text. if this process is performed using a large body of text, the successor variety of substrings of a term will decrease as more characters are added until a segment limit is obtained. At this level, the successor variety will consequently increase. This kind of information will be used to identify stems. We can also apply the same process for any other language, for example in Arabic language, if we have the following body of the text: (علم، عتيق، عائق، علو، عمق، علوم، عوام، علوج، علومه). The successor varieties of the word "علومهم" will be as follows: ("ع", (05), ("عل", (01), ("علوم", (02), ("علو", (02)).

*Table lookup method:* Some stemming algorithms use a table or a list of index terms and their stems stored in advance to extract the stem of a new word that occurs on the index. Therefore, terms from queries and indexes could then be stemmed based on this table lookup [9] using some well-known techniques such as B-tree or Hash table. In this case, such lookups would be very fast. For example, "used", "users", "using", and "useful". All these terms can be stemmed to the common stem "use". with this approach, three problems can be met. The first is that to build these lookup tables we need to work directly on a language and thus require prior linguistic knowledge related to this language. The second is that these tables may miss out on some exceptional cases. A third problem is the storage cost for such a table (the required memory space).

Table 1. Principle of Table Lookup Method

| Term          | Stem   |
|---------------|--------|
| English       |        |
| Information   | Inform |
| Informatics   | Inform |
| Inform        | Inform |
| Informational | Inform |
| Arabic        |        |
| علوم          | علم    |
| عالم          | علم    |
| معلم          | علم    |
| علماء         | علم    |
| معلمة         | علم    |
| معلومات       | علم    |
| علماء         | علم    |
| يتعلمون       | علم    |

*N-Grams method*: Proposed by [10] and also called the shared di-gram method. A di-gram (a bigram) is a pair of consecutive letters, but we can also use 3-grams, 4-grams, ..., and hence it's called the n-grams method. With this method, two words are associated based on the common unique di-grams they both share. To calculate this association measure, we generally use Dice's coefficient [11]. For example, the following terms; policy, political can be segmented into bigrams as follows:

policy => po ol li ic cy (05 digrams)  
 unique digrams = po ol li ic cy (05 digrams)  
 political => po ol li it ti ic ca al (08 digrams)  
 unique digrams => po ol li it ti ic ca al (08 digrams)  
 common unique digrams => po ol li ic (04 digrams)

Thus, "policy" has (05) digrams, all are unique, and "political" has (08) digrams, all are unique. The two words share (04) unique bigrams.

Once the common unique bigrams for the word pair have been identified and counted, a similarity metric based on them is computed. The first used similarity metric was Dice's coefficient, which is defined as follows: [11]

$$S = 2 * C / (A + B) \quad (1)$$

where A and B are the numbers of unique bigrams in the first and the second words successively, C the number of unique bigrams shared by A and B.

For the example above, Dice's coefficient would equal:  $(2 \times 4) / (5 + 8) = 8 / 13 = 0.61$ . Such similarity metrics are determined for all pairs of terms in the database. Once such similarity metric is computed for all the word pairs they are clustered as groups.

### 3. Literature Review: Stemming Algorithms

Stemming algorithms can be classified into three main groups: truncating methods, statistical methods, and mixed methods (hybrid methods) [12]. Each group has its way to find the stem of the word.

#### A. Truncating Methods

This group of methods is related to removing the affixes of a word. One well-known stemmer in this group is the Lovins stemmer proposed by Lovins in 1968. It performs a lookup on a table that contains some endings, conditions, and transformation rules [13]. The Lovins stemmer always proceeds to the removal of the longest suffix from a word. And hence, the

word is recoded using a separate table that makes various adjustments to convert these stems into correct words. As advantages of this algorithm, we note that it is very fast and can handle the removal of double letters in words like ‘setting’ being transformed to ‘set’ and also handles many irregular plurals like – appendix and appendices, etc. The drawbacks of the Lovins approach are that it is time and data-consuming. Furthermore, it is possible that many suffixes are not available in the table of endings.

Porter stemming algorithm [14, 15] is one of the best and the most popular stemming methods until our days, proposed in 1980 by Porter. Many improvements have been introduced to the basic algorithm. The main idea of this stemmer is that the suffixes in English are often made up of a combination of smaller and simpler suffixes. It works through many steps, within each step, many rules are applied until one of them verifies the required conditions. If the selected rule is accepted, the suffix is removed accordingly, and the stemming process passes to the next step until achieving the last step in which the resultant stem is returned. Porter has also built a detailed framework of stemming which is known as the ‘Snowball stemming framework’. This framework allows programmers to develop their stemmers for different languages. The Paice/Husk stemmer is an iterative algorithm based on a table containing a large list of rules (about 120 rules) indexed by the last letter of a suffix [6]. On each iteration, it tries to find a valid rule according to the last character of the word. Each rule performs one of two possible tasks, a deletion or replacement of an ending. If there is no such rule, the algorithm terminates. As advantages of Porter stemmer, we note: it is a simple stemmer, and every iteration performs both deletion and replacement according to the applied rule. The disadvantage is it is a very heavy algorithm and over stemming may occur in some cases. Dawson Stemmer [16]: This stemmer is an extension of the Lovins algorithm except that it covers a much more comprehensive list of suffixes (about 1200 suffixes). Like Lovins, it is a single pass stemmer and hence it is relatively fast. The suffixes are indexed by their length and last letter and then stored according to reversed order. Therefore, they are organized as a set of character trees which permit a rapid access. The advantage is that it covers more suffixes than Lovins stemmer and it is fast in execution. The disadvantage is it is very complex and does not have a unique standard implementation.

### *B. Statistical Methods*

This group of methods is called statistical methods, it contains stemmers that are based on statistical analysis and techniques. Most of the methods remove also the affixes after implementing some statistical procedures. In this group we can find the following stemmers: N-Grams Stemmer [6, 17]: it is considered as a language-independent stemmer in which a string-similarity approach is used to convert word inflation to its stem. An n-gram can be defined as a set of  $n$  successive characters extracted from a word. The basic idea of this approach is that similar words share a high proportion of n-grams. For  $n$  equals 2, the words extracted are called digrams. For example, the word ‘information’ results in the generation of the digrams: \*i, in, nf, fo, or, rm, ma, at, ti, io, on, n\* Where ‘\*’ denotes a padding space. usually,  $n$  can take a value of 4 or 5. The major advantage of this stemmer is that it is completely language independent and hence very useful in a large range of applications. as disadvantage is it is memory space-consuming when creating and storing these n-grams. HMM Stemmer: This stemmer was proposed by [18] and is mainly based on the concept of the Hidden Markov Model (HMMs) which is a finite-state automaton. At each transition, the new state provides a character with a calculated probability. This method does not need a prior linguistic knowledge of the dataset, and the probability to take each path can be computed and the most probable path is found using the Viterbi coding in the associated automata graph. YASS stemmer (Yet Another Stripping Stemmer) was proposed by [19]. It is considered as a statistical and corpus-based stemmer at the same time since it does not rely on prior linguistic knowledge. It seems effective for languages that are suffixing in nature, and its performance is comparable to of some well-known stemmers such as Porter and Lovins in terms of average precision and the total number of relevant retrieved documents. In YASS approach, a set of

boolean string distance measures is defined [20], this distance measure is used to check the similarity between two words A and B by mapping them to a real number R, where a smaller value of R indicates a greater similarity between A and B.

### *C. Mixed Methods*

This group groups a variety of mixed methods which are:

**The Inflectional and Derivational Methods:** Require a deep morphologic analysis and a very large corpus to develop these stemmers, and hence they are part of corpus base stemmers. In inflectional methods, the word variants are mainly related to the variation of language syntax like plural, gender, case, etc. Whereas, in derivational methods, the word variants are related only to the part-of-speech (POS) of a sentence in which the word occurs [12].

**Krovetz Stemmer (KSTEM):** was presented in 1993 by [21], it is a linguistic lexical stemmer, very complicated, and it effectively removes inflectional suffixes in three steps. Unfortunately, Krovetz stemmer is not able to find the stems for all word variants, perhaps it is the reason for which this stemmer is generally used as a pre-stemmer before applying any other stemming algorithm namely: Porter, Lovins, which can increase the speed and the effectiveness of the main stemmer. Krovetz stemmer is also considered as a light stemmer and does not give a good recall and precision performance.

**Xerox Inflectional and Derivational Analyzer:** many years ago, the linguistics groups at Xerox corporation have developed a lexical database for English and some other languages, this lexical database helps in analyzing and generating inflectional and derivational morphology of words. It reduces each surface word to the form which can be found in the dictionary, as follows [22]: nouns singular (e.g. children child), verbs infinitive (e.g. understood understand), etc. The advantages of this stemmer are that it works well with a large document and also removes the prefixes where ever applicable. The obtained stems are valid since a lexical database providing a morphological analysis of any word in the lexicon is available. Among the drawbacks of this stemmer is that the output depends mainly on the quality of lexical database which may not be exhaustive. Another drawback is that this method is based on a lexicon, hence, it cannot always give the correct stems which does not appear in the lexicon.

**Corpus Based Stemmer** This method of stemming was proposed by [17]. This method tries to overcome some of the drawbacks of the Porter stemmer. The corpus-based stemming refers to an automatic modification of conflation classes – words that have resulted in a common stem, to suit the characteristics of a given text corpus using statistical methods. The main hypothesis is that word forms that should be conflated for a predefined corpus will co-occur in documents of the same corpus. Using this concept, many errors of over stemming and under stemming can be resolved e.g. It is also important to note that this stemmer works within two steps, in the first step, it uses the Porter stemmer to identify the stems of conflated words and in the second step, it uses the corpus statistics to redefine the conflation. As advantage of this stemmer that it can potentially avoid making conflations that are not appropriate for a given corpus and the result is an actual word not an incomplete stem. The disadvantage is that you need to develop the statistical measure for every corpus separately and the processing time increases as in the first step, and stemming algorithms are first used before using this method.

**The Context Sensitive Stemmer** is done using statistical modeling on the query side. This method was proposed by [23]. The basic idea behind this stemmer is that before submitting any query to the search engine, we predict morphological variants which would be useful in the search process. Using this method, the number of bad expansions can be reduced, which in turn reduces the cost of additional computation and consequently improves the precision at the same time. After the derivation of the predicted word variants from the query, a context-sensitive document matching is done for these variants. The major advantage of this stemmer is that it improves selective word expansion on the query side and matches conservative word occurrence on the document side. Two important disadvantages can arise here, the processing time and the complexity of the algorithm.

Many Stemming algorithms have been developed for a wide range of languages including English [24], Latin [22], Swedish [25], German and Italian [26], French [27], Turkish [28, 29]. For Indian languages [30-32], we can note: Punjabi [31] with an accuracy of 80.73%, Bengali [33] gives an accuracy of 96.27%, Urdu [32] based on two approaches: Length based and Frequency-based, gives for each approach successively one of the following accuracy: 84.27% and 79.63%, Hindi [30] gives an accuracy of 91.59%, Another Hindi stemmer was proposed in 2014 by [34], based on Suffix Stripping of porter and gives an accuracy of 83.65%. For other non-Indian languages, we can note: Persian, and Indonesian [35, 31, 32].

For the Arabic Language, there are three different Stemming approaches: the root-based approach [36-41]; the light stemmer approach [42-46]; and the statistical stemmer approach (N-Grams) [47-49]. But, until now, no perfect stemmer for this language is available. Figure 2 summarizes the different stemming algorithms.

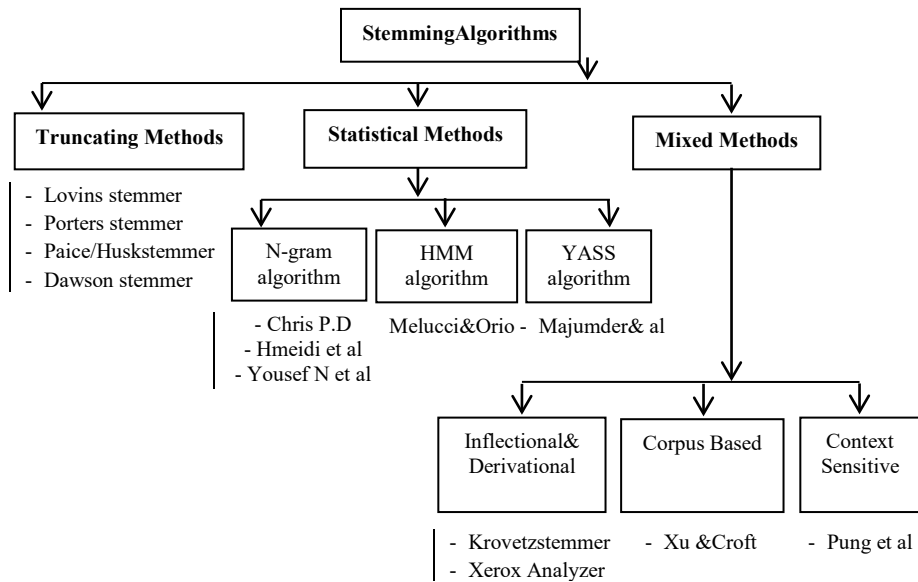


Figure 2. Classification of stemming Algorithms [12]

#### 4. Our Stemming Approach

Our stemming approach is characterized by: the use of the n-grams of characters technique and the extraction of the word root [50-53]. In the experimental part, we have applied this approach on three different languages which are: Arabic, English, and French. The stemming process can be divided into many phases:

Phase 1: segmentation of the word subject of stemming, and all the candidate roots into bigrams (2-grams). Table 2 describes the segmentation phase for three different words given in the three languages presented above.

Phase 2: Computation of stemming parameters which are:

$N_w$  : The size of the word  $W$  in number of bigrams.

$N_{R_i}$  : The size of the root  $R_i$  in number of bigrams  $N_{wR_i}$  : The number of common bigrams between the word  $W$  and the root  $R_i$

$N_{w\bar{R}_i}$  : The number of bigrams belonging to the word  $W$  and do not belong to the root  $R_i$

$$(N_{w\bar{R}_i} = N_w - N_{wR_i})$$

$N_{R_i\bar{W}}$  : The number of bigrams belonging to the root  $R_i$  and do not belong to the word  $W(N_{R_i\bar{W}} = N_{R_i} - N_{wR_i})$ .

Table 2. Description of the segmentation phase (Phase 1)

| Language | Word (W)/Bigrams                                    | List of roots (R <sub>i</sub> )/Bigrams                                                                                                                                                                                                                                                                                                                                                                                  |
|----------|-----------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Arabic   | يفقدون<br>يف، فق، قد، دو، ون                        | $R_1 = \text{"نَجَح"} \rightarrow (\text{نج، جح})$<br>$R_2 = \text{"خَرَج"} \rightarrow (\text{خر، رج})$<br>$R_3 = \text{"فَقَد"} \rightarrow (\text{فق، قد})$<br>$R_4 = \text{"عَقَد"} \rightarrow (\text{عق، قد})$                                                                                                                                                                                                     |
| French   | Connaissance<br>co on nn na ai is ss sa<br>an nc ce | $R_1 = \text{"command"} \rightarrow (\text{co, om, mm, ma, an, nd})$<br>$R_2 = \text{"avanc"} \rightarrow (\text{av, va, an, nc})$<br>$R_3 = \text{"concev"} \rightarrow (\text{co, on, nc, ce, ev})$<br>$R_4 = \text{"commenc"} \rightarrow (\text{co, om, mm, me, en, nc})$<br>$R_5 = \text{"connait"} \rightarrow (\text{co, on, nn, na, ai, it})$<br>$R_6 = \text{"conclu"} \rightarrow (\text{co, on, nc, cl, lu})$ |
| English  | Transaction<br>tr ra an ns sa ac ct ti io<br>on     | $R_1 = \text{"action"} \rightarrow (\text{ac, ct, ti, io, on})$<br>$R_2 = \text{"extra"} \rightarrow (\text{ex, xt, tr, ra})$<br>$R_3 = \text{"intra"} \rightarrow (\text{in, nt, tr, ra})$<br>$R_4 = \text{"dict"} \rightarrow (\text{di, ic, ct})$                                                                                                                                                                     |

For the previous example we have the results shown in Table 3 below:

Table 3. Computation of stemming parameters (phase 2)

| Word(W)      | $N_w$ | Associated roots $R_i$                                                                                                                               | $N_{R_i}$            | $N_{wR_i}$           | $N_{w\bar{R}_i}$      | $N_{R_i\bar{W}}$     |
|--------------|-------|------------------------------------------------------------------------------------------------------------------------------------------------------|----------------------|----------------------|-----------------------|----------------------|
| يفقدون       | 05    | $R_1 = \text{"نَجَح"}, R_2 = \text{"خَرَج"}$<br>$R_3 = \text{"فَقَد"}, R_4 = \text{"عَقَد"}$                                                         | 2, 2<br>2, 2         | 0, 0<br>2, 1         | 5, 5<br>3, 4          | 2, 2<br>0, 1         |
| connaissance | 11    | $R_1 = \text{"command"}, R_2 = \text{"avanc"}$<br>$R_3 = \text{"concev"}, R_4 = \text{"commenc"}$<br>$R_5 = \text{"connait"}, R_6 = \text{"conclu"}$ | 6, 4<br>5, 6<br>6, 5 | 2, 0<br>4, 2<br>5, 3 | 9, 11<br>7, 9<br>6, 8 | 4, 4<br>1, 4<br>1, 2 |
| transaction  | 10    | $R_1 = \text{"action"}, R_2 = \text{"extra"}$<br>$R_3 = \text{"intra"}, R_4 = \text{"dict"}$                                                         | 5, 4<br>4, 3         | 5, 2<br>2, 1         | 5, 8<br>8, 9          | 0, 2<br>2, 2         |

Phase 3: Selection of candidate roots which are those sharing at least one bigram with the word  $W$  ( $N_{wR_i} \geq 1$ ) among the predefined list of roots. In our previous example, we have obtained the results stored in Table 4:

Table 4. Selection of candidate roots (phase 3)

| Word(W)      | $N_w$ | Associated roots $R_i$                                                                                                           | $N_{R_i}$         | $N_{wR_i}$        | $N_{w\bar{R}_i}$  | $N_{R_i\bar{W}}$  |
|--------------|-------|----------------------------------------------------------------------------------------------------------------------------------|-------------------|-------------------|-------------------|-------------------|
| يفقدون       | 05    | $R_3 = \text{"فَقَد"}, R_4 = \text{"عَقَد"}$                                                                                     | 2, 2              | 2, 1              | 3, 4              | 0, 1              |
| connaissance | 11    | $R_1 = \text{"command"}, R_3 = \text{"concev"},$<br>$R_4 = \text{"commenc"}, R_5 = \text{"connait"},$<br>$R_6 = \text{"conclu"}$ | 6, 5<br>6, 6<br>5 | 2, 4<br>2, 5<br>3 | 9, 7<br>9, 6<br>8 | 4, 1<br>4, 1<br>2 |
| transaction  | 10    | $R_1 = \text{"action"}, R_2 = \text{"extra"}$<br>$R_3 = \text{"intra"}, R_4 = \text{"dict"}$                                     | 5, 4<br>4, 3      | 5, 2<br>2, 1      | 5, 8<br>8, 9      | 0, 2<br>2, 2      |

Phase 4: Computation of the distance  $D(W, R_i)$  between the word  $W$  and each candidate root  $R_i$  according to the formula (4):

$$D(W, R_i) = \frac{N_{w\bar{R}_i} + N_{R_i\bar{w}}}{N_w + N_{R_i}} \quad (2)$$

In the case of our previous example, we obtain the results shown in Table 5 as follows:

Table 5. Computation of the distance  $D(W, R_i)$

| Word( $W$ )  | $N_w$ | Associated roots $R_i$                                                                                                                 | $N_{R_i}$         | $N_{wR_i}$        | $N_{w\bar{R}_i}$  | $N_{R_i\bar{w}}$  | Dist.Val<br>$D(W, R_i)$                      |
|--------------|-------|----------------------------------------------------------------------------------------------------------------------------------------|-------------------|-------------------|-------------------|-------------------|----------------------------------------------|
| يفقدون       | 05    | $R_3 = \text{"فقد"} , R_4 = \text{"عقد"} "$                                                                                            | 2, 2              | 2, 1              | 3, 4              | 0, 1              | 0.4285,<br>0.7142                            |
| connaissance | 11    | $R_1 = \text{"command"} , R_3 = \text{"concev"} ,$<br>$R_4 = \text{"commenc"} , R_5 = \text{"connaît"} , R_6$<br>$= \text{"conclu"} "$ | 6, 5<br>6, 6<br>5 | 2, 4<br>2, 5<br>3 | 9, 7<br>9, 6<br>8 | 4, 1<br>4, 1<br>2 | 0,7647,<br>0,5<br>0,7647,<br>0,4117<br>0,625 |
| transaction  | 10    | $R_1 = \text{"action"} , R_2 = \text{"extra"} "$<br>$R_3 = \text{"intra"} , R_4 = \text{"dict"} "$                                     | 5, 4<br>4, 3      | 5, 2<br>2, 1      | 5, 8<br>8, 9      | 0, 2<br>2, 2      | 0.3333,<br>0.7142<br>0.7142,<br>0.8461       |

Phase 5: assigning the root having the minimum value of distance  $D(W, R_i)$  to the word  $W$ .

In the previous example, the roots assigned to the given words are illustrated on Table 6:

Table 6. Extraction of the word root (step 5)

| Word ( $W$ ) | Extracted root( $R$ ) | Effective root |
|--------------|-----------------------|----------------|
| يفقدون       | فقد                   | فقد            |
| connaissance | connaît               | connaît        |
| transaction  | action                | action         |

## 5. Experimental Results and Interpretation

To validate the new proposed stemming approach, we have performed many experiments on the three languages cited previously. These experiments were as follows:

### A. Used Dataset

We have used for each language three corpora with different sizes: small corpus, middle corpus, and large corpus (see Table 7). In each there exist three files as described below:

1. The file of derived forms (gross words) that contains several morphological forms of words derived from a list of roots.
2. The file of roots containing a collection of roots, we note here that for the Arabic language, these roots may be: 3-lateral, 4-lateral, 5-lateral, and 6-lateral (see Table 8). We underline also that some of them are vocalic roots that contain at least one vowel.
3. The file of golden roots containing the list of effective roots which may be matched to the words present in each corpus, this golden list has been established by an expert linguist and used as a reference list, i.e., the calculation of the root extraction accuracy (success ratio) is done by comparison between the list of obtained roots (extracted by the system) and the reference list (established by the expert). The extraction process and the obtained results are shown in Tables 9, 10, 11 as well as figures 3, 4, 5, 6, 7, 8.

Table 7. The Used Dataset

| Corpus       | Language | Size of the derived words file | Size of the roots file | Size of the golden roots file |
|--------------|----------|--------------------------------|------------------------|-------------------------------|
| Small corpus | Arabic   | 50                             | 25                     | 50                            |
|              | French   | 44                             | 36                     | 44                            |
|              | English  | 92                             | 56                     | 92                            |
| Middle       | Arabic   | 300                            | 150                    | 300                           |

| Corpus       | Language | Size of the derived words file | Size of the roots file | Size of the golden roots file |
|--------------|----------|--------------------------------|------------------------|-------------------------------|
| corpus       | French   | 250                            | 220                    | 250                           |
|              | English  | 450                            | 180                    | 450                           |
| Large Corpus | Arabic   | 2000                           | 650                    | 2000                          |
|              | French   | 1500                           | 600                    | 1500                          |
|              | English  | 1500                           | 620                    | 1500                          |

Table 8. Examples of Arabic roots (3-lateral, 4-lateral, 5-lateral, 6-lateral)

| Trilateral roots | Quadrilateral roots | Quinquelateral roots | Hexalateral roots |
|------------------|---------------------|----------------------|-------------------|
| ز ر ع            | أ ك ر م             | ا ن ط ل ق            | ا س ت ع م ل       |
| ص ن ع            | أ ع ا ن             | ا ن ك س ر            | ا س ت ح س ن       |
| ج م ع            | ح ط م               | ا ق ت ص د            | ا خ ش و ش ن       |
| أ ب ي            | ع ل م               | ا ج ت م ع            | ا ع ش و ش ب       |
| س ع ل            | ط م ا ن             | ت ن ا ز ل            | ا ق ش ع ر         |

### B. Obtained Results

After applying our proposed stemming approach on the three languages: Arabic, French, and English, we obtained the results presented in the tables and figures below:

Table 9. Illustration of the obtained results after the segmentation phase

| Word(W)        | N-grams                                | Nb.Ng<br>( $N_W$ ) | Root    | N-grams           | Nb.Ng<br>( $N_{R_i}$ ) |
|----------------|----------------------------------------|--------------------|---------|-------------------|------------------------|
| <i>Arabic</i>  |                                        |                    |         |                   |                        |
| يتعلمون        | ون يت تع عل ل م                        | 7                  | علم     | عل ل م            | 3                      |
| اقتصاد         | اق قت تص صا اد                         | 5                  | اقتصد   | اق قت تص ص        | 4                      |
| سنستدرجهم      | سن نس ست تد در رج جه هم                | 8                  | درج     | در رج             | 2                      |
| <i>French</i>  |                                        |                    |         |                   |                        |
| apprentissage  | ap pp pr re en nt ti is ss<br>sa ag ge | 12                 | apprend | ap pp pr re en nd | 6                      |
| calculatrice   | ca al lc cu ul la at tr ri ic<br>ce    | 11                 | calcul  | ca al lc cu ul    | 5                      |
| connaissance   | co on nn na ai is ss sa an<br>nc ce    | 11                 | connaît | co on nn na aî ît | 6                      |
| <i>English</i> |                                        |                    |         |                   |                        |
| bicycle        | bi ic cy yc cl le                      | 6                  | cycle   | Cy yc cl le       | 4                      |
| transactions   | tr ra an ns sa ac ct ti io<br>on ns    | 11                 | action  | ac ct ti io on    | 5                      |
| geopolitics    | Ge eo op po ol li it ti ic<br>cs       | 10                 | geo     | Ge eo             | 2                      |

Table 10. Extraction of word roots using the new stemming approach.

| Word (W)      | Nearest Roots $R_i$      | Nb.Common Bigrams<br>( $N_{WR_i}$ ) | Distance Values<br>$D(W, R_i)$    | Extracted Root | Correct Root |
|---------------|--------------------------|-------------------------------------|-----------------------------------|----------------|--------------|
| <i>Arabic</i> |                          |                                     |                                   |                |              |
| يتعلمون       | كلم ، علم ، علم ،<br>علج | 1 ، 1 ، 3 ، 2                       | 0.77 , 0.77 ,<br><u>0.4</u> , 0.6 | علم            | علم          |
| سنستدرجهم     | درج                      | 2                                   | <u>0.6</u>                        | درج            | درج          |
| <i>French</i> |                          |                                     |                                   |                |              |
| Apprentissage | Apprend,                 | 4, 1, 2, 1, 2,                      | <u>0.55</u> , 0.9, 0.78,          | apprend        | apprend      |

| Word ( $W$ )   | Nearest Roots $R_i$                                          | Nb.Common Bigrams ( $N_{WR_i}$ ) | Distance Values $D(W, R_i)$                      | Extracted Root | Correct Root |
|----------------|--------------------------------------------------------------|----------------------------------|--------------------------------------------------|----------------|--------------|
|                | assembl, assist,<br>associ, assur,<br>autoris,<br>comprend   | 1, 2                             | 0.89, 0.77, 0.9,<br>0.78                         |                |              |
| connaissance   | avanc,<br>command,<br>commenc,<br>concev, conclu,<br>connaît | 3, 2, 3, 4, 3,<br>4              | 0.64, 0.78,<br>0.68, <u>0.57</u> ,<br>0.66, 0.57 | concev         | connaît      |
| <i>English</i> |                                                              |                                  |                                                  |                |              |
| transactions   | action, extra,<br>intra, dict                                | 5, 2, 2, 1                       | <u>0.28</u> , 0.69,<br>0.69, 0.83                | action         | action       |
| geopolitics    | Action, geo,<br>poly                                         | 1, 2, 2                          | 0.86, <u>0.66</u> , 0.69                         | geo            | geo          |

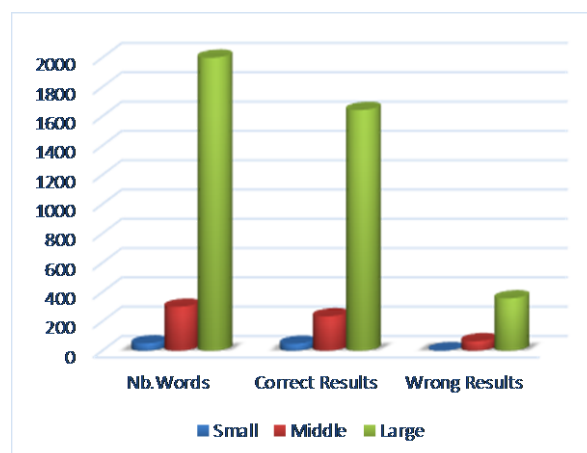


Figure 3. Illustration of the obtained results in terms of number of words (Arabic)

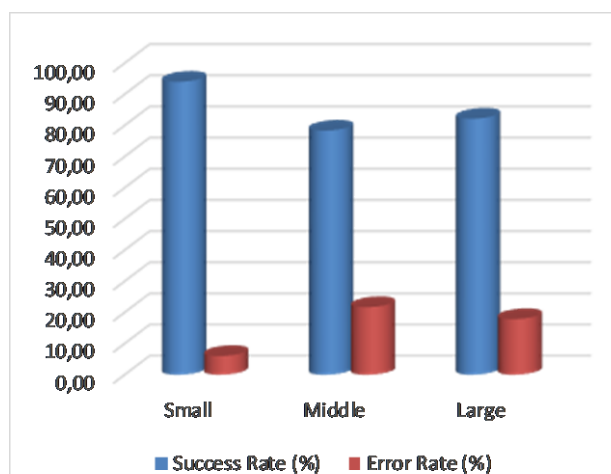


Figure 4. Computation of success and error rates (Arabic)

Table 11. Extraction of word roots: Global results.

| Corpus        | Language | Nb. Roots | Nb. Words | Correct Results | Wrong Results | Success Rate (%) | Error Rate (%) |
|---------------|----------|-----------|-----------|-----------------|---------------|------------------|----------------|
| Small Corpus  | Ar       | 25        | 50        | 47              | 03            | 94,00            | 06,00          |
|               | Fr       | 36        | 44        | 41              | 03            | 93,18            | 06,82          |
|               | En       | 56        | 92        | 85              | 07            | 92,39            | 07,61          |
| Middle Corpus | Ar       | 150       | 300       | 235             | 65            | 78,33            | 21,67          |
|               | Fr       | 220       | 250       | 217             | 33            | 86,80            | 13,20          |
|               | En       | 180       | 450       | 411             | 39            | 91,33            | 08,67          |
| Large Corpus  | Ar       | 650       | 2000      | 1643            | 357           | 82,15            | 17,85          |
|               | Fr       | 600       | 1500      | 1173            | 327           | 78,20            | 21,80          |
|               | En       | 620       | 1500      | 1327            | 173           | 88,46            | 11,54          |

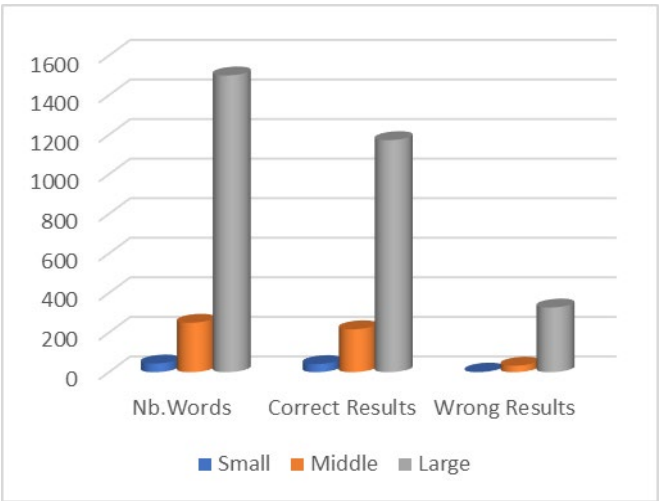


Figure 5. Illustration of the obtained results in terms of number of words (French)

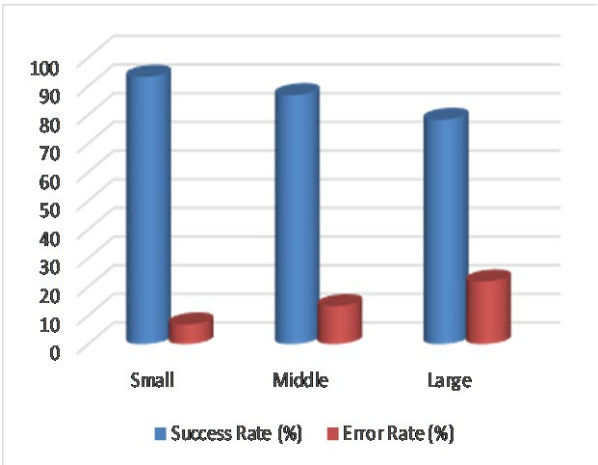


Figure 6. Computation of success and error rates (French)

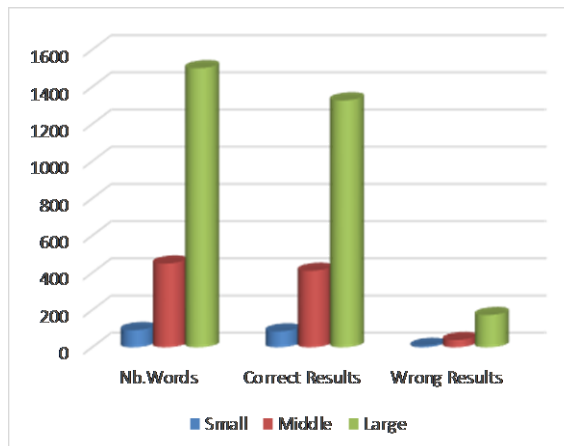


Figure 7. Illustration of the obtained results in terms of number of words (English)

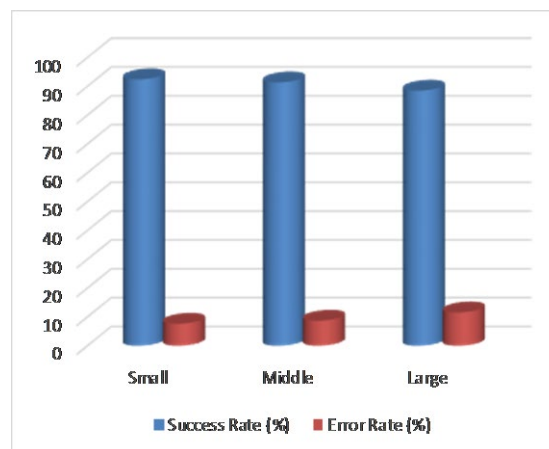


Figure 8. Computation of success and error rates (English)

## 6. Advantages of our Stemming Approach

Our new stemming approach has many advantages among them we can note the following:

1. It is language-independent, so, we can apply it to many languages at the same time and without modification (a multilingual approach), Which is not the case for the existing stemmers, for example, the Porter stemming framework “snowball” allows the researchers in different languages to develop their stemmers (section 3), but they must modify the basic algorithm to deal with the morphology and grammatical rules of each language (a kind of language dependence).
2. Does not require any affix removal, where there is a real risk to remove some native letters of the word instead of the affixes.
3. Valid for any length of Arabic root. (e.g., 3-lateral, 4-lateral, 5-lateral, and 6-lateral roots).
4. Works very well for vocalic and strong roots, these two classes of roots usually pose problems in Arabic during their derivation, since they change their forms completely. We note here that no Arabic stemmer covers this property.
5. It is mainly based on a simple calculation of distances (it is a full statistical stemmer). Thus, no morphological rules or grammatical patterns are used during the extraction process.
6. It is a practical stemmer and then easy to implement.
7. Can be extended to the lemmatization process.

## 7. Discussion

In the first order, it seems important to underline that it was more interesting to establish a comparison between our multilingual stemming approach and some existing ones, but this was not possible, for a simple reason that all available stemmers are monolingual, i.e., each language has a set of stemmers based on a distinct theoretical background and different degrees of performances. This lets to conclude, that the major advantage of our stemming approach is that it is more general and completely language independent (Multilingual). On the other hand, we see in table 11, that the performance of the stemming algorithm in our approach decreases when the words base increases. This can be interpreted by the fact that the increase of the words' base must be also followed by the increase of the roots' base too. If not, many words will not be associated with their roots in the golden list (reference list of roots). A second remark that must be mentioned here, is the low value of accuracy related to Arabic stemming compared with English and French, this is due to the method used to segment words and roots into bigrams of characters. For the method we have used in the present work, a bigram is obtained by concatenation of two direct successive letters. This will deal with Latin languages such as English and French for which the root is a sequence of successive letters extracted from the beginning, the middle, or the end of the word itself. E.g., for English we can meet pairs (word, root) as follows: (Geography, Geo), and the same for French (Confirmation, Confirm). This is completely different and will not deal with Arabic whose root must be built with a series of letters that cannot mandatorily be consecutive, and consequently the number of common n-grams between the word and each candidate root decreases as well as the distance between them. It is the major reason for which it will be a large divergence between a word and its candidate roots. As it is mentioned in the following examples: (علماء, علم) - (كُتَّاب, كتب), where roots are formed of no consecutive letters. to overcome this problem, and so improve the stemming accuracy for Arabic, we suggest the following segmentation method: to form sequences of bigrams when segmenting a word or a root, we take all possible combinations of characters not only the pairs of successive ones: (e.g., علماء □ علم, ماء □ ماء, ل □ لاء, لا □ لا, ع □ عاء, عم □ علم).

## 8. Comparison with Other Existing Stemmers

In this section, we establish a comparison of our proposed stemmer with some well-known stemmers as follows:

### A. Comparison according to their Characteristics

In table 12, we give a large comparison between existing stemmers in terms of: used approach, main advantages and disadvantages.

Table 12. Comparison with existing stemmers according to their characteristics

| Stemmer            | Type/Used approach                  | Advantages                                                                                                      | Disadvantages                                                                                                    |
|--------------------|-------------------------------------|-----------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------|
| Lovins stemmer     | Truncating method (Affixes removal) | -Fast stemmer<br>- Simple pass<br>- Removal of double letters in words<br>- Handles many irregular plural words | - Time consuming<br>- Frequently fails to form words from stems<br>- Depends on the used vocabulary              |
| Porter stemmer     | Truncating method (Affixes removal) | -Produces best results of stemming<br>- Less error rates<br>- It is a light stemmer                             | - Generated stems are not always the real words.<br>-Works on several steps (05 steps).<br>- Consumes more time. |
| Paice/Husk stemmer | Truncating method (Affixes removal) | A simple stemmer<br>Each iteration considers                                                                    | Heavy algorithm<br>Over stemming is                                                                              |

| Stemmer                          | Type/Used approach                                | Advantages                                                                                                                                                                                                                                                                                                                                                          | Disadvantages                                                                                                                   |
|----------------------------------|---------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------|
|                                  |                                                   | both removal and replacement                                                                                                                                                                                                                                                                                                                                        | possible                                                                                                                        |
| Dawson stemmer                   | Truncating method (Affixes removal)               | -Covers a large list of suffixes (1200)<br>-Fast in execution.                                                                                                                                                                                                                                                                                                      | - enough complex<br>- no standard implementation                                                                                |
| Krovetz stemmer                  | Inflectional & derivational method                | - It is a light stemmer.<br>- Can be used as a pre-stemmer for other stemmers.                                                                                                                                                                                                                                                                                      | -Not valid for documents of large sizes<br>-Enable to take care of words outside the lexicon<br>-does not produce good results. |
| Youssef N et al stemmer (Arabic) | Statistical method (N-grams approach)             | - Based on bigrams and string comparison<br>- does not require morphological rules<br>- Language-independent                                                                                                                                                                                                                                                        | - Space-consuming<br>-Time-consuming.<br>-monolingual stemmer (Arabic).                                                         |
| Our proposed stemmer             | Statistical method (An improved n-grams approach) | - Multilingual stemmer.<br>-Based on an improved n-grams approach<br>- does not require morphological rules<br>- Language-independent<br>- For Arabic, it is valid for any length of the root. (3, 4, 5, 6-lateral roots).<br>- works well for strong and vocalic roots.<br>- it is practical and then easy to implement.<br>- It can be applied for lemmatization. | -Time-consuming for a big corpus.<br>-Based on a reference base of roots.                                                       |

### *B. Comparison according to the obtained Accuracy*

As it was mentioned in section 7 (Discussion), it seems important to underline that it was more interesting to establish a comparison between our multilingual stemming approach and some existing ones, but this was not possible, for a simple reason that all available stemmers are monolingual, i.e., each language has a set of stemmers based on a distinct theoretical background and different degrees of performances. This let's to conclude, that the major advantage of our stemming approach is that it is more general and completely language independent (Multilingual). But, in this section, we try to establish a small comparison between our stemmer with other monolingual stemmers in terms of the obtained accuracy. For this purpose, we consider our multilingual stemmer as 03 independent stemmers (Arabic stemmer, French stemmer, and English stemmer). We take also the average value of the accuracy which combine the obtained values for the three used corpora (small, middle, large). Table 13 and Figure 9 bellow summarize the established comparison.

Table 13. Comparison with existing stemmers according to the obtained accuracy

| <i>Stemmer</i>              | <i>Type/Language</i>             | <i>Obtained Accuracy</i>                                                |
|-----------------------------|----------------------------------|-------------------------------------------------------------------------|
| <i>English [24]</i>         | <i>Monolingual (English)</i>     | 90.00%                                                                  |
| <i>Turish [28]</i>          | <i>Monolingual (Turkish)</i>     | 93.83%                                                                  |
| <i>Khodja [37]</i>          | <i>Monolingual (Arabic)</i>      | 68.00%                                                                  |
| <i>Yousef et al [49]</i>    | <i>Monolingual (Arabic)</i>      | 92.00%                                                                  |
| <i>Punjabi [31]</i>         | <i>Monolingual (Punjabi)</i>     | 80.73%                                                                  |
| <i>Bengali [33]</i>         | <i>Monolingual (Bengali)</i>     | 96.27%                                                                  |
| <i>Urdo [32]</i>            | <i>Monolingual (Urdo)</i>        | 84.27%                                                                  |
| <i>Hindi [30]</i>           | <i>Monolingual (Hindi)</i>       | 91.95%                                                                  |
| <i>Hindi [31]</i>           | <i>Monolingual (Hindi)</i>       | 83.65%                                                                  |
| <i>Our Proposed Stemmer</i> | <i>Multilingual (Ar, En, Fr)</i> | Ar → avg. acc: 84.82%<br>En → avg. acc: 90.72%<br>Fr → avg. acc: 86.06% |

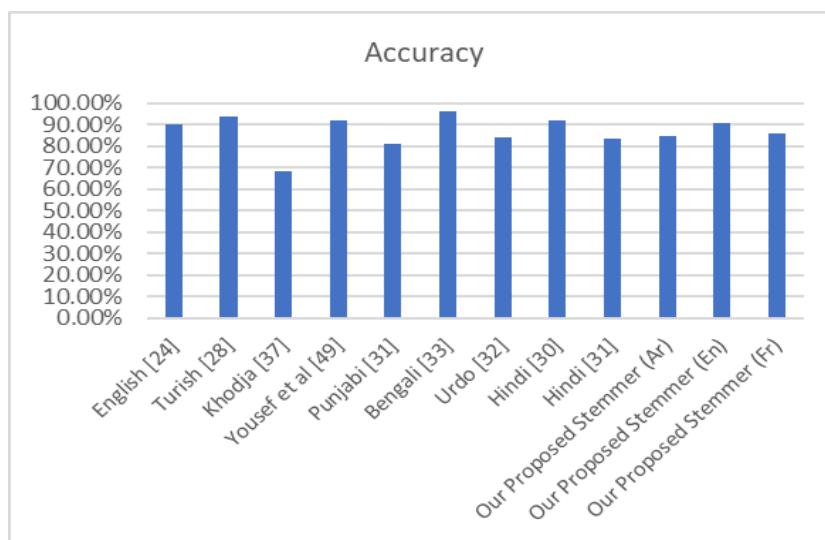


Figure 9. Comparison with existing stemmers according to the obtained accuracy

## 9. Conclusion and Future Suggestions

In the present paper, we have studied one of the most critical problems in linguistics and NLP areas which is the problem of stemming and its direct influence on the quality of TC, IR, and many other fields of research such as NLP and text mining tasks. This linguistic technique can reduce the size of terms vocabulary by 30 to 40% in TC, and increase the effectiveness of IRS. Throughout our article, we presented the well-known approaches, namely: morphological methods (truncating methods), statistical methods, and mixed methods that combine the two previous ones. For each class of methods, we have exposed briefly the main advantages and disadvantages. In the light of the presented methods, we proposed a multilingual stemming approach, which is characterized by: being completely language independent, based on a statistical approach that in turn uses the technique of n-grams, does not require any prior linguistic knowledge, able to find all kinds of roots (strong and vocalic), which is very perfect for Arabic since it is considered as one of the richest languages, with complex morphology and a difficult structure. The realized experiments prove that the obtained accuracy when extracting the root for the studied languages is very satisfactory and can be improved in other future projects in the same field. Some of the improvements that we can propose are as follows:

- Extend the present algorithm to lemmatization.

- Take into consideration other languages such as German, Spanish, Turkish, etc.
- Apply our algorithm to other big datasets.
- Use another non-successive segmentation method for Arabic to improve the extraction of the root as it was explained in the discussion section.
- Improve the success rate for the three languages by introducing semantic processing.
- Use the developed stemmer to increase the quality of multilingual categorization in TC, and the effectiveness multilingual search in IR

## 10. Acknowledgements

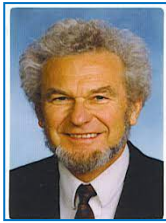
I thank in the first order the research team of the University of Vienna (Wien Universitat), Austria at which I have done a short internship during the period 1st, October – 10th, October 2018, especially Pr. Dr. Wolfgang Klas the head of CS Department, Dr. Bilal AbuNaim, Ms. Elahieh M-O, Ms. Manuella the secretary of the department, For their good reception and very important advice during the period of my internship. I can't also forget all people who help me to achieve this work.

## 11. References

- [1]. Mesleh, A.W. SVM based Arabic Language Text Classification System: Feature Selection Comparative study. *In the 12th WSEAS International Conference on applied mathematics* (pp. 31-29.), Cairo, Egypt, 2007
- [2]. Ionan, P. Web Document Classification and its Performance Evaluation. *In the 9th WSEAS International Conference on evolutionary computing EC08* (pp.353-348), Sofia, Bulgaria, May 2-4, 2008
- [3]. Komkid, C., Narodom, K., Kittisak, K., & Netaya, K. Comparison of Feature Selection and Classification Algorithms for Restaurant Dataset Classification. *In the 11th conference on Latest Advances in Systems Science and Computational Intelligence* (pp.134-129), 2012
- [4]. Hadni, M., Ouattik, S.A., & Lachkar, A. Effective Arabic stemmer-based hybrid approach for Arabic text categorization. *International Journal of Data Mining and Knowledge Management Process (IJDKP)*, 3(4), 14-1, 2013
- [5]. Chris, P.D. Another stemmer. *ACM SIGIR*, 24(3), 61-56, 1990
- [6]. Chris, P.D. An evaluation method for stemming algorithms. *In the 17th annual International ACM SIGIR conference on Research and development in information retrieval* (pp.50-42), 1994
- [7]. Deepika, S. Stemming Algorithms: A Comparative Study and their Analysis. *International Journal of Applied Information Systems (IJ AIS)*, 4(3), 12-7, 2012
- [8]. Hafer, M., & Weiss, S. Word Segmentation by Letter Successor Varieties. *Information Storage and Retrieval*, 10, 371-86, 1974
- [9]. Frakes, W.B. Term Conation for Information Retrieval: In Research and Development. *In Information Retrieval*. ed. C. van Rijsbergen, New York: Cambridge University Press, 1984
- [10]. Adamson, G., Boreham, G. The Use of an Association Measure Based on Character Structure to Identify Semantically related pairs of words and document titles. *Information Storage and Retrieval*, 10(1), 253-60, 1974
- [11]. Frakes, W.B. Stemming Algorithm in Information Retrieval Data Structures and Algorithm ed. C. van Rijsbergen, Chapter 8: 139-132, 1992
- [12]. Jivani, A. A Comparative Study of Stemming Algorithms. *International Journal of Computer Technology and Applications IJCTA*, 2(6), 1938-1930, 2011
- [13]. Lovins, J.B. Development of a stemming algorithm. *Mechanical Translation and computational Linguistics*, 11(2), 31-22, 1968
- [14]. Porter, M.F. An algorithm for suffix stripping. *Program*, 14(3), 137-130, 1980
- [15]. Porter, M.F. Snowball: A language for stemming Algorithm, October 2001.

- [16]. Dawson, J. Suffix removal and word conation. *Bulletin of the Association for Literary and Linguistic Computing*, 2(3), 47-33, 1974
- [17]. Croft, B.W., Xu, J. Corpus-based stemming using co-occurrence of word variants. *ACM Transactions on Information Systems*, 16(1), 81-61, 1998
- [18]. Melucci, M., Orio, N. A novel method for stemmer generation based on hidden Markov models. In the twelfth inter. conf on Info & knowledge management (pp. 138-131), 2003
- [19]. Majumder, M., Mitra, S., Parui, K., Kole, G., & Mitra, P. YASS: Yet another Suffix Stripper. *ACM Transaction on Information Systems (TOIS)*, 25(4),5-6, 2007
- [20]. Levenstein, V.I. Binary codes capable of correcting deletions, insertions and reversals. *Soviet Physics Doklady*, 10(8), 710-707, 1966
- [21]. Krovetz, R. Viewing morphology as an inference process. *In the 16th annual International ACM SIGIR conference on Research and development in information retrieval* (pp.202-191), 1993
- [22]. Hull, D.A., Grefenstette, A. A detailed analysis of English Stemming Algorithms. XEROX Technical Report on <http://www.xrce.xerox>, 1996
- [23]. Funchun, P., Nawaas, A., Xin, L. and Yumao, L. Context sensitive stemming for web search. *In the 30th annual international ACM SIGIR conference on Research and development in IR* (pp.646-639), 2007
- [24]. Greengrass, M., Robertson, A.M., Robyn, S., & Willett, P. Context sensitive stemming for web search in processing morphological variants in searches of Latin text. *Information research news*; 6(4), 5-2, 1996
- [25]. Carlberger, J., Dalianis, H., Hassel, M., & Knutsson, O. Improving precision in Information retrieval for Swedish using stemming. *In The 3th Nordic conference on computational linguistics NODALIDA'01* (pp.22-21), Uppsala, Sweden, 2001.
- [26]. Monz, C., & Rijke, M. Shallow morphological analysis in monolingual information retrieval for Dutch, German and Italian. *In the 2nd workshop of Cross-language information retrieval and evaluation Form CLEF 2001* (pp.359-357), C. Peters Ed, Springer Verlag, 2001
- [27]. Moulinier, I., McCulloh, A., & Lund, E. Non-English monolingual retrieval. *In Cross-language information retrieval and evaluation. In the CLEF 2000 workshop* (pp.187-176), C. Peters Ed, Springer Verlag, 2001
- [28]. Ekmekcioglu, F.C., Lynch, M.F., & Willett, P. Stemming and N-Gram matching for term conation in Turkish texts. *Information Research News*, 7(1),6-2, 1996
- [29]. Tan, A.H., & Yu, P. A Comparative Study on Chinese Text Categorization Methods. *workshop on Text and Web Mining* (pp.35-24), Melbourne, 2000
- [30]. Thangarasu, M., & Manavalan, R. A Literature Review: Stemming Algorithms for Indian Languages. *International Journal of Computer Trends and Technology (IJCTT)*, 4(8), 2584-2582, 2013
- [31]. diBijal, D., & Sanket, S. Overview of Stemming Algorithm for Indian and Non-Indian Languages. *International Journal of Computer Sciences and Information Technologies (IJCSIT)*, 5(2), 1146-1144, 2014
- [32]. Kasthuri, M., & Kumar, S.B.R. A comprehensive analyze of stemming algorithms for Indian and non-indian languages. *Inte.Jo.Comp.Eng.and.App IJCA*, 7(3), 8-1, 2014
- [33]. Suprabhat, D., & Pabitra, M. A Rule-based Approach of Stemming for Inflectional and Derivational Words in Bengali. *In the IEEE Students' Technology Symposium* (pp.136-134), jan 14-16, 2001, Kharagpur, 2011
- [34]. Gupta, V. Hindi Rule Based Stemmer for Nouns. *International Journal of Advanced Research in Computer Science and Software Engineering*; 4(1), 65-62, 2014
- [35]. Husain, M.S. An unsupervised approach to develop stemmer. *International Journal on Natural Language Computing IJNLC*, 1(2),23-15, 2012
- [36]. Al-Fedaghi, S., & Al-Anzi, F. A new algorithm to generate Arabic root-pattern forms. *In the 11th national Computer Conference and Exhibition* (pp. 400-391), 1989

- [37]. Khodja, S. Stemming Arabic Text. Technical Report, Computing Department, Lancaster University, 1999
- [38]. Larkey, L., & Connell, M.E. Arabic information retrieval at UMass in TREC-10. *In the 10th retrieval conference TREC-2001* (pp.570-562), Gaithersburg: NIST, 2002
- [39]. Momani, M.; & Faraj, J. A novel algorithm to extract tri- literal Arabic roots, New Rules to Directly Extract Roots of Arabic Words. *In the International Conference on Computer Systems and Applications* (pp.315-309). *IEEE Xplore Press*, 2009
- [40]. Al-Nashashibi, M.Y., Neagu, D., & Yaghi, A.A. An improved root Extraction technique for arabic words. *In 2nd International Conference on Computer Technology and Development ICCTD* (pp.269-264). *IEEE Xplore Press*, 2010
- [41]. AbuHawas, F., & Keith, E. Rule- based Approach for Arabic Root Extraction: New Rules to Directly Extract Roots of Arabic Words. *Journal of Computing and Information Technology - CIT*, 22(1),68-57, 2014
- [42]. Al-shalabi, R., Kanaan, G., & Al-Serhan, H. New Approach for Extracting Arabic Roots. *In the International Arab Conference on Information Technology ACIT03* (pp.59-42). Alexandria, Egypt, 2003
- [43]. Kanaan, G., Al-Shalabi, R., & Al-Kabi, M. New approach for Extracting quadrilateral Arabic roots. *Journal of Basic Science and Engineering*, 14(1),66-51, 2005
- [44]. Al-Nashashibi, M.Y., Neagu, D., & Yaghi, A.A. Stemming techniques for Arabic words: A comparative study. *In 2nd International Conference on Computer Technology and development CCTD* (pp.276-270), 2010
- [45]. Mohamad, A., Al-Shalabi, R., Kanaan, G., & Al-Nobani, A. Building an effective Rule-Based light stemmer for Arabic language to improve search effectiveness. *The International Arab Journal of Information Technology IAJIT*, 9(4),372-368, 2012
- [46]. Al-omari, A., Abuata, B., & Al-kabi, M. Building and Benchmarking New Heavy-Light Arabic Stemmer. *In 4th International Conference on Information and Communication Systems ICICS* (pp.69-63), 2013
- [47]. Mustafa, S.H., & Al-Radaideh, Q.A. Using n-grams for Arabic text searching. *Journal of the American Society for Information Science and Technology*; 55, 1007-1002, 2004
- [48]. Hmeidi, I., Al-Shalabi, R., Al-Taani, A.T., Najadat, H., & Al-Hazaimeh, S.A. A novel approach to the extraction of roots from Arabic words using bigrams. *Journal of the Association for Information Science and Technology*, 61(3),591-583, 2010
- [49]. Yousef, N., Aymen, A.E., Ashraf, O., & Hayel, K. An Improved Arabic Words Roots Extraction Method Using N-gram Technique. *Journal of Computer science JSC*, 10(4), 716, 2014
- [50]. Gadri, S., Moussaoui, A.O. Information Retrieval: A New Multilingual Stemmer Based on a Statistical Approach. *In 3th International conference on Control, Engineering & Information Technology (CEIT'15)*, Tlemcen, Algeria, 25-27 may 2015.
- [51]. Gadri, S., Moussaoui, A.O. Arabic Text Categorization: An Improved Algorithm Based on Ngrams technique to Extract Arabic Words Roots, *International Arab Journal of Information Technology IAJIT*, online Version Vol.14, No.06, November 2017. ([http://ccis2k.org/iajit/?option=com\\_content&task=blogcategory&id=125&Itemid=427](http://ccis2k.org/iajit/?option=com_content&task=blogcategory&id=125&Itemid=427)).
- [52]. Gadri, S., Moussaoui, A.O. Arabic Information Retrieval: Influence of Stemming on the Effectiveness of Search in IRS, *7th International Conference on advanced Technology ICAT'18*, April, 28 – May, 1st, 2018, Antalya, Turkey
- [53]. Gadri, S., Moussaoui, A.O. Multilingual Text Categorization: Increasing the Quality of Categorization by a Statistical Stemming Approach, *International Conference on intelligent Information Processing Security and Advanced Communication (IPAC 2015)*, November 23-25, 2015, ACM Conference. Batna, Algeria.



**Said Gadri**, Received his degree of engineer in computer science, field Hard&Software from the University Ferhat Abbas of Setif, Algeria in 1996, the degree of magister (Bac + 7) from the University Mohamed Boudiaf of M'sila, Algeria in 2006, and the degree of doctor of sciences from the university Ferhat Abbas of Setif, Algeria in 2016, the degree of habilitation from the university Mohamed Boudiaf of M'sila in 2021. He is currently an associate professor (permanent teacher) of computer science at the University Mohamed Boudiaf of M'sila, Algeria since 2007. Teaching different subjects like Graph theory, Computer architecture, Databases, Web technology, programming languages (Pascal, C++, C#, Algorithmic), Expert systems, Modeling and simulation, Artificial Intelligence, Information Retrieval, etc. He is a member of the scientific council of mathematics and computer sciences faculty, since 2009 to our days, a member of the teaching committee of the CS department (2013 – 2015). Currently, he is interested in many areas of research such as Text categorization, machine learning, Text mining, Information retrieval, Deep Learning, and natural language processing. Published more than 30 papers in different international journals and conferences around the world.



**Erich Neuhold** is a professor at the faculty of computer science at the University of Vienna and the faculty of computer science at the Darmstadt University of Technology until 2005. He was also Director of the Fraunhofer institute for integrated publication and information systems (IPSI) in Darmstadt. Earlier he has been a professor at the university of Stuttgart and the technical University of Vienna and he has also worked in research and management positions for IBM and Hewlett Packard both in Europe and the USA. His area of expertise include: distributed databases, object-oriented databases, databases for internet (e.g., unstructured, semi-structured and structured), information retrieval, information visualization, workflow and business processes and their applications in digital libraries, cultural heritage, e-science, e-commerce, and e-government. He has published 04 books and about 200 papers. His work has appeared, among others, in the VLDB journal, information systems, Acta informatica and many conferences as, for example, VLDB, ICDE, MMDB, ADL, DL, IRC, CaiSE, EC-Web, EURASIA, etc. He has served in all capacities on many conference committees and was a PC Chair and a General Chair among others, of VLDB, ICDE, WISE, SAINT, JCDL and ECDL. He was the chair of the IEEE TCDE and IEEE TCDL and was the chair of the ICDE steering committee and he is currently the chair of the JCDL steering committee. He is a fellow of IEEE, USA and the Gesellschaft für Informatik, Germany.