

Semantic Elements Extraction based on Syntactic Structure of Arabic Sentences

Raghda Hraiz, Mariam Khader, Arafat Awajan, Akram Alkous
Computer Science
Princess Sumaya University for Technology
Amman, Jordan

Abstract— Huge amount of data is uploaded to the Internet daily, most of this data are a user-generated data. User-generated data impose challenges in processing by machines because of its unstructured nature. This paper proposes a new approach for extracting semantic relations from Arabic sentences. The proposed approach analyzes the syntactic structure of Arabic sentences and recognizes its semantic elements. The proposed approach consists of three stages: syntactic analysis, segmentation and semantic relation extraction phase. The extracted elements from the text are organized in XML format. The proposed approach has been evaluated using a dataset of 15 sentences from AL-Jazeera website. The results have shown that the proposed method is reliable and promising in extracting semantic elements.

Keywords—*Semantic Network; Semantic Elements; Information Extraction; Natural Language Processing; knowledge Representation.*

I. INTRODUCTION

Huge amount of data is uploaded to the internet every day; most of the data is generated on social networks [1], these data are mainly unstructured (natural language text). Many researches were proposed in order to enable machines to process unstructured data; those researches included data representation using logic, rules, frames, etc [2]. One the most implemented representation is the semantic networks. The semantic network is defined as a graphical representation of knowledge, where concepts represented as nodes and their relations as arcs. Semantic networks were used to represent knowledge and to infer conclusions from this knowledge [3].

This work proposes a new approach for semantic elements extraction between entities in Arabic texts using semantic networks. The extracted elements are generated from the syntactic structure of the sentences. The results show that the proposed method is promising for Arabic semantic, as it works automatically and does not require identifying rules at the semantic level, which in general require identifying a huge number of semantic patterns and rules. Instead of that, the method relies on the syntactic level to define patterns and rules. Which require defining less number of rules, and can work on any text from any domain.

The rest of the paper is organized as follow; next section presents the state of the related work in extracting semantic relations for both English and Arabic texts. Section III, details the methodology steps. In section IV, evaluation of the proposed approach is demonstrated. Finally, section V concludes the result of applying syntactic structure for semantic elements extraction on Arabic sentences.

II. RELATED WORK

Semantic Network has been presented early at 1966 by Ross Quillian. It was proposed first as a general mechanism for representing knowledge [4]. The semantic network has been used in natural language processing as a knowledge representation in order to represent the semantic relations between concepts. It is form of directed or undirected graph, where vertices represent concepts, and edges represent semantic relations between those concepts [1].

Different approaches have been applied in order to extract entities and semantic relations between those entities in user-generated text. The most common methods found in literature are pattern-based methods. A pattern is a template or a structure of a sentence that describes how related words could occur in it. Those patterns can be prepared manually by an expert or generated automatically. In [5], Aldhubayi & Alyahya have described three main approaches that are used to generate semantic patterns automatically. Those approaches are supervised machine learning methods, semi-supervised pattern-based methods and bootstrapping methods. Supervised machine learning methods extract semantic relations by applying classification techniques on the dataset. For using classification, a dataset of positive and negative test cases for semantic relations is needed. However, preparing adequate annotated test cases is expensive task and requires huge efforts and human expertise. On the other hand, semi-supervised and bootstrapping methods require small datasets of patterns of defined semantic relations to begin the extraction process.

An unsupervised Relation Extraction method was proposed by [6]. The developed system of this method discovered entities and relations from Wikipedia documents. The benefit of using Wikipedia was to utilize the clean and well-written structure and contents of Wikipedia documents, especially the



English content. The system of relation extraction included three steps: text preprocessing and concept pair collecting, web context collecting and dependency pattern extracting and clustering Algorithm. In text preprocessing and concept pair collecting step, the method scans Wikipedia documents in order to extract concept pairs. Wikipedia was utilized because its document contains "anchor-text", which connects the current article with its related articles. This method considers the current document and any of its related documents as a concept pairs and extracts any sentence containing both concepts together. In web context collecting step, the system will query Google with the concept pair; the result will be processed to extract information about the relations between the concept pair. This step was important in order to generate well-defined and accurate relations. In the last step, an entropy-based algorithm was employed to compare and rank all possible terms, those terms were used to represent the relations. The terms are all nouns and verbs extracted from previous steps. In order to improve the quality of the relation, functional terms were added by the system to the relation to make it readable.

In other researches, relation extraction method was developed for specific domains by [1]. In their research, a method for discovering medical elements between entities was proposed. In order to recognize syntactic entities in the medical documents, the researchers proposed an enhancement on an existing Medical Entities recognition system called MetaMap [7]. The research included identifying the category of each extracted name entity such as health problem, medical treatment or medical procedure and then extracting the semantic relations between extracted entities. The proposed system defined the types of the semantic relations. There are six main types: treats, prevents, causes, complicates, diagnoses and sign or symptom of. In order to extract semantic relations, the authors designed several linguistic patterns for each semantic relation type. The pattern was determined based on the selected sentences from PubMed Central articles. The sentences were selected using UMLS Metathesaurus knowledge and MeSH indexing of PubMed Central. The linguistic patterns were compared with each sentence in the corpus in order to choose the most suitable relation. The proposed method had good results on a corpus consist of 580 sentences. The same authors also presented MeTAE, which is a platform for extracting information from medical texts.

In their research Bollegala et al. [8] proposed extensional and intentional representations of the semantic relations between entity pairs. The extensional representation of semantic relation was achieved by deriving all possible entity pairs where this semantic relation holds. In the other hand, the intentional representation was achieved by identifying a set of lexical-syntactic patterns. The syntactic were used to construct the semantic relations. In order to extract the potential entities, the authors proposed an efficient single-pass extraction method. The proposed method scans all the sentences in the

corpus once and work as a part of speech tagger and a noun-phrase chunker in order to extract entities. After that, the method apply a subsequence pattern-mining algorithm, which discovers lexical syntactic patterns. The discovered patterns were used to represent the semantic relation between two entities. Finally, the lexical patterns were clustered using a sequential co-clustering algorithm. Results of the patterns clustering had promising results on the ENT benchmark dataset. The experiments proved the applicability of the proposed method on large dataset. The method was able to classify 53 relation types correctly in a large online social network system, with more than 35 million edges.

KNOWITALL is an autonomous system that was proposed by [9]. This system was used to scan web pages and identify concepts, relationships and facts. Specifically, KNOWITALL relied on an expandable ontology and a set of rules. The system has a domain and languageindependent architecture to fill the ontology with new relations and facts. KNOWITALL was developed to be scalable, where the modules ran as threads and the modules communication achieved by asynchronous message passing.

In [10], WOE (Wikipedia-based Open Extractor) system was proposed. WOE is an open information extractor system. It uses self-supervised learning to generate training samples, and then uses those training samples to train an open information extractor. The training samples are generated by matching the attributes in Infobox (in Wikipedia pages) with the corresponding sentences in the natural text.

Regarding Arabic language, there were few studies for semantic relation extraction. The reason is that there are not enough resources to serve such researches. In the research of Lahbib et al. [11], the authors suggested a hybrid method for extracting semantic relation from Arabic documents. The suggested method incorporated statistical techniques and linguistic knowledge for discovering semantic relations. The method studied the relation between syntactic dependencies and the semantic of the sentences. The focus was on noun phrases, specifically, the following types of noun phrases: adjective phrases, prepositional phrases, annexation phrases, conjunctive phrase and complex noun phrases. After discovering the relations between syntactic dependencies and the semantic elements, the method utilized statistical techniques to measure the similarity between semantic entities. In order to evaluate the accuracy of the proposed method, the authors tested it on vocalized text (Al Hadith) to reduce ambiguity.

Other researches on Arabic focused on extracting special types of relations. In their research [5] the goal of the research was on extracting antonym relationships between nouns. The authors designed a pattern-based bootstrapping approach based on Arabic language corpora called arTenTen and a corpus analysis tool called Sketch Engine. The method focused on discovering the most frequent patterns surrounds antonym pairs

in Arabic language. The evaluation results showed the ability of this method to discover antonym pairs with a precision of 76%.

Despite the importance of semantic relation extraction research, there are limited work and the result are not satisfactory yet. In this paper, a new approach has been presented. The proposed approach exploits the syntactic structure of the sentences in Arabic text to extract the semantic relations. The work will focus on ten different structures to study the relation between syntactic structure and semantic elements. There are several structures templates for Arabic sentences, however for the purpose of this work; most ten general structures were used as a template as it will be illustrated in next section.

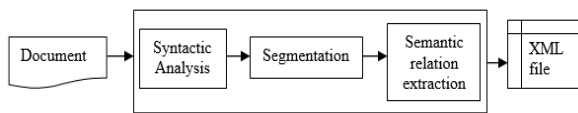


Figure 1: Main Architecture of the proposed system

III. METHODOLOGY

This paper proposes an approach for extracting semantic elements based on the relation between syntactic structure and semantic elements of Arabic sentences. Arabic is a Semitic language that has a rich and complicated syntax structure. In order to study and analyze its syntax, a dedicated linguistic science called “Al-Naho” was created, and this science helps to reveal the semantic of the sentences [12].

According to “Al-Naho” science, the sentences in Arabic may start with a noun (Noun phrase) or start with a verb (verb phrase), where each type may have several orders. In this research, a sample of ten templates (patterns) is analyzed in order to study the ability of extracting semantic element (entities and relations) based on syntactic structure. In order to apply the approach, two preprocessing stages should be done on the selected documents. Besides that, lexicons of synonyms and name entity should be prepared. In this work, the preprocessing stages and lexicons have been prepared manually. In addition, a system was developed in order to extract those elements automatically. The input of the system is a document of natural language text and the output is an XML file that contains the extracted entities and relations. Figure 1 describes the different elements of the proposed system, those elements are: the input document, preprocessing steps, semantic relation extraction and the XML output file.

The system consists of two preprocessing stages, which are explained in next two subsections: syntactic analysis, and segmentation.

The first preprocessing stage is syntactic analysis. In this stage, each word in the input document has been analyzed according to the rules of “Al-Naho”; those rules were extracted from “Al-Rajhi” (الراجحي) book. In this stage, each word is

assigned a tag to represent its syntactic role. Table 1 shows a sample of the tags that have been used in this research. Determining these tags is very important for semantic analysis, because in Arabic, the syntactic tags are needed to determine the specific type of the word science.

Table 1: sample of syntactic tags used in this research

Syntactic role	Tag	Meaning
مبتدأ	SubNP	Subject in Noun phrase
خبر	Pred	Predicate
نعت	Adj	Adjective
فاعل	SubVP	Subject in verbal phrase
ظرف زمان	TmAdv	Temporal adverb
ظرف مكان	DirAdv	Directional adverbs
بدل	Cnt	Counterpart
مفعول به	Obj	Object
فعل	V	Verb
اسم مجرور	Gen	Genitive
تمييز	Disc	Discrimination
حال	STT	State
حرف جر	Prep	Preposition
حرف عطف	Conn	Conjunctive adverbs

For example, each noun in Arabic has specific type; the noun could be one of the following (and many other tags) [12]:

- Subject, which might describe the person/concept who did the event specified by the verb
- Object, which might describe the person/concept who has affected by the event
- "مبتدأ" (Subject in Noun phrase) which is the first word in noun phrases, and the most important noun in the sentence
- "خبر" (predicate) which give information about the "مبتدأ" (Subject in Noun phrase)
- "حال" (State), which describe the guise of the subject when the event happened
- "نعت" (Adjective), which is a property of the previous noun in the sentence
- "تمييز" (discrimination), which describe the previous mentioned noun or verb, used to mention its type and eliminate ambiguity

A. Segmentation

It is a very hard to deal with the whole document while extracting semantic relations. In this phase, the input documents are segmented into sentences. The segmentation process in this stage is done based on syntactic structure of the sentence. A set of words is considered a sentence if it form a complete sentence from syntactic perspective, for example the sentence “مات العالم” (the scientist died) is considered a complete sentence because it consist of a verb and a subject and the verb is Intransitive verb. Whereas the sentence “اعلم خالد اياه” (Khalid informed his father) is not a complete sentence, because the

verb "اعلم" is Transitive verb that requires two objects according to the Arabic syntax rules [12].

Extracting semantic elements

The proposed semantic extraction approach work as follow, first, the system parse the document and process each sentence individually by comparing each sentence with a set of defined rules (which will be discussed later). If the sentence matches the rule, the semantic elements will be extracted according to that rule. Finally, the semantic elements will be saved in XML file.

The two lexicons needed for the purpose of this work are a name entity lexicon and synonyms lexicon. The name entity lexicon is required to provide the semantic meaning of the sentence. For example, name entity is important to recognize that "Syria" is a country not a person name and "Microsoft" is an organization. In addition, name entity is important when multiple words are used to represent the same entity such as "Hashemite Kingdom of Jordan" and "Jordan". The second lexicon contains the synonyms of the words. This lexicon is important when several words have the same meaning. For example, the words "تحيات" means live, so instead of recognizing them as two different relations, the system will check the synonyms lexicon and if the two words are synonyms, the system will consider them the same relation.

The XML file consists of three main elements, entities, relations and attributes. Entity is an independent concept that represents person, organization, country, etc. Relation is a link that connects two related entities together, this link should have name, and may have place and time. The attribute is a property of entity or relation that adds extra information about it. The main difference between the attribute and the entity is that the attribute is always related to an entity. The following is an example of the semantic elements:

"الف الشافعي العالم المشهور كتاب الرسالة في مصر"

"The well-known scientist Al-Shafee wrote Al-Resalah book in Egypt"

In the above example, there are three entities (Al-Shafee, Al-Resalah Book and Egypt). The first entity has two attribute (well-known and scientist). There is only one relation in this sentence, which is wrote. Wrote relation connects Al-Shafee to Al-Resalah Book. This relation will be added to the first entity (Al-Shafee), which will be used later to connect it to the second entity. The extracted relation also contains place and time, the time in this example is not stated so it will be empty, and the place is Egypt.

To extract the semantic elements a set of rules has been defined, which are discussed next. Each rule consists of two parts: syntactic template and processing steps, where each template is used to determine the processing required on the tested sentence. A syntactic template is a sequence of tags that are used to define the syntactic structure of the sentence. Each

sentence will be compared with the syntactic templates, if a sentence matches a specific syntactic template, the rule corresponding to the matched template will be fired. If a rule is fired, the preprocessing steps according to the syntactic template are applied to extract the semantic elements from the sentence.

The following illustrate the analysis of the ten rules corresponding to the syntactic templates, the tags are according to table 1:

- Template 1 $\rightarrow$ {"SubNP", "Pred"}

If the sentence contains two words tagged as "SubNP" followed by "Pred", the first word will be added as an entity to the semantic graph. The second word will have two choices; the first choice is that if there is an entity that represents the same concept, a new relation will be added to the entity of the first word and will be connected to the second one. The second choice is that if the second word does not represent an entity (does not have its own attributes or relations), the word will be considered as an attribute for the first entity.

- Template 2 $\rightarrow$ {"V", "SubVP", "Obj"}

If the sentence contains three words tagged as "V" followed by "SubVP" and then "Obj", the system will search its pool for the second and third words. If any of these words does not exist, the system will add it as entity. Then it will create a relation in the entity of the second word, and connect it to the third word. The name of the relation in this case will be the first word.

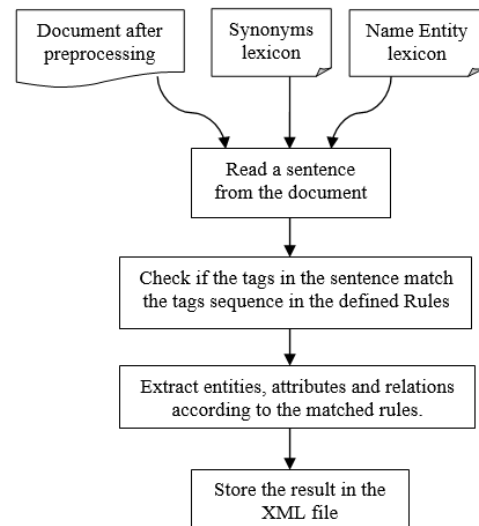


Figure 2: Main Steps of the proposed system after preprocessing

- Template 3 $\rightarrow$ {"V", "Sub", "Prep", "Gen"}

If the sentence contains the following sequence of tags "V", "Sub", "Prep" and "Gen". The system will search for the second and forth words. If any of the two words does not exist, the system will add it as an entity. Then the system will create a relation in the entity of the second word, and connect it to the fourth word. The name of the relation will be the connection of the first word and the third word.

- Template 4 → {"SubNP", "V", "Prep", "Gen"}

*this rule will be fired if the word that has the "Gen" tag is found in Name Entity lexicon as place.

In case the sentence contains the following sequence of tags "SubNP", "V", "Prep" and "Gen", the system will search for the first and forth words. If the words not exist, the system will add them as entities. Then it will create a relation in the entity of the first word that connects it to the fourth word. The name of the relation will be the concatenation of the second word and the third word.

- Template 5 → {"SubNP", "Adj"}

If the sentence contains two words tagged as "SubNP" followed by "Adj", the first word will be added as an entity to the semantic graph. The second word will have two choices; the first choice is that if there is an entity that represents the same concept, a new relation will be added to the entity of the first word and this added relation will be used to connect the first word to the second one. The second choice is that if the word does not represent an entity (does not have its own attributes or relations), the word will be considered as an attribute of the first entity.

- Template 6 → {"V", "Sub", "Obj", "TmAdv"}

If the sentence contains the following sequence of tags "V", "Sub", "Obj" and "TmAdv", the system will search for the second and fourth words. If the words are not exist, the system will add them as entities. After that, the system will create a relation in the entity of the second word; this relation will be used to connect the second word to the fourth word. The name of the relation will be the first word in the sentence. A new attribute will be added to the relation, the name of the new attribute will be the third word.

- Template 7 → {"V", "SubVP", "Obj", "TmAdv"}

If the sentence consists of the following sequence of tags "V", "SubVP", "Obj" and "TmAdv", the system will search for the second, third and fourth words. If the words do not exist, the system will add them as entities. Then the system will create a relation in the entity of the second word; this relation will be used to connect second word to the third word. The name of the relation will be the first word. The fourth word will be added as the time of the relation.

- Template 8 → {"Pred", "Adj"}

If the sentence consists of two words tagged as "Pred" followed by "Adj", the system will search for the first word. If it does not exist, it will be added as an entity to the semantic graph. The second word will have two choices; the first choice is that if there is an entity that represents the same concept, a new relation will be added to the entity of the first word and will be used to connect the first word to the second one. The second choice is that if the word does not represent an entity (does not have its own attributes or relations), the second word will be considered as an attribute of the first entity.

- Template 9 → {"V", "SubVP", "Obj", "Stt"}

In case the sentence contains the following sequence of tag "V", "SubVP", "Obj" and "Stt", the system will search for the second and third words. If the words do not exist, the system will add them as entities. Then the system will create a relation in the entity of the second word, this relation will be used to connect the second word to the third word. The name of the relation will be the first word. A new attribute will be added to the relation. The name of the new attribute will be the fourth word.

- Template 10 → {"V", "SubVP", "Obj", "DirAdv"}

If the sentence contains the following sequence of tags "V", "SubVP", "Obj" and "DirAdv", the system will search for the second, third and fourth words. If the words do not exist, the system will add them as entities. Then the system will create a relation in the entity of the second word, this relation will be used to connect the second word to the third word. The name of the relation will be the first word. Then the fourth word will be the place of the relation.

IV. EVALUATION

In order to evaluate the reliability of the proposed approach, a set of 15 sentences have been chosen from Aljazera website as a data set [13]. The used test data set is shown in table 2. The output of the system has been compared with human expert analysis results. According to the human analysis, the dataset contains 14 entities, 37 relation and 14 attributes. The proposed system has correctly identified 13 entities, 9 relations and 9 attributes. The system was not able to recognize all the semantic elements. The reason for misidentifying the elements was that the used templates in the system do not cover all the syntactic templates of Arabic sentences. In addition, the system failed when the sentence contained multi-word expressions such as "رضي الله عنه". The system has also failed when the text contained similes or metaphors. Additionally, the current representation of the rules needs to be improved in a way that is capable to handle larger number of templates. The rule representation can be improved by creating special type of finite state transducer that can parse the sentence and generate the output efficiently.

The following two examples show how the system performed in semantic elements extraction. Example one

shows a sentence that have been analyzed correctly by the system. The system creates an entity tag in the xml output for the second and fourth word according to template number three. According to the templates rules, the system created a relation that connects the entity of the second word with the entity of the fourth word. The name of the relation is the first and third words as mentioned in rule number three.

Example 1:

Gen_جامعة ديوبند الإسلامية:Prep_في:SubVP_البنوري:V_تخرج

Meaning:

AL-Banoree graduated from Dayobond Islamic university

Output:

```
<Entity name="البنوري">
  <Relation connected_to="جامعة ديوبند الإسلامية" name="تخرج في"/>
</Entity>
<Entity name="جامعة ديوبند الإسلامية"/>
```

Meaning:

```
<Entity name="AL - Banoree">
  <Relation connected_to="Dayobond Islamic university"
  name="graduated from"/>
</Entity>
<Entity name="Dayobond Islamic university"/>
```

Example two shows that the system has not succeed in extracting semantic elements. The reason was that the input sentence has not matched any template in the system.

Example 2:

Conn.:الصلاح:Conn.:ب:SubVP_مديرية مردان الهندية:V_تشتهر

Meaning:

Indian Mardan Directorate is known by its Science and righteousness

V. CONCLUSION

In this paper, a new method has been proposed to extract semantic relations from Arabic text. The proposed approach exploits the syntactic structure of the sentence to determine the semantic elements. The method consists of three phases, Syntactic analysis, segmentation and semantic relation extracting. The results show that the proposed method is promising for Arabic semantic, as it works automatically and does not require identifying rules at the semantic level, which in general require identifying a huge number of semantic patterns and rules. Instead of that, the method relies on the syntactic level to define patterns and rules. Which require defining less number of rules, and can work on any text from any domain.

Table2 : The used test data set for semantic elements extraction

Num.	Sentence
------	----------

1	محمد يوسف البنوري عالم هندي متضلّع في التفسير والحديث ومن أبرز سدنة اللغة العربية.
2	تبوأ مناصب علمية رفيعة في الهند وباكستان،
3	وآلف كتباً قيمة في عدة مجالات.
4	ولد محمد يوسف بن محمد زكريا البنوري عام 1908 في قرية مهايت اباد التابعة لمديرية مردان الهندية، لأسرة يتصل نسبها بعلي بن أبي رضي الله عنه
5	وتشتهر بالعلم والصلاح
6	درس البنوري العلوم الدينية على أبيه وخاله الشيخ فضل حمداني في ببشاور
7	ودرس النحو والصرف في كابل بأفغانستان على يد عبد القادر اللمقاني ومحمد صالح الغليغوري،
8	كما تلقى العلم على الشيخ شبير أحمد ومحمد أنور شاه الكشميري وغيرهما.
9	وفي عام 1929 تخرج في جامعة ديوبند الإسلامية العريقة في الهند.
10	عمل البنوري استاذاً في الجامعة الإسلامية ببادهل في مقاطعة مومباي الهندية،
11	وترأس جمعية علماء الهند، إلى جانب مسؤوليات ووظائف دينية في باكستان
12	تتميز التجربة العلمية للبنوري بتضلّعه في تفسير القرآن ومعرفة بالحديث الشريف،
13	ويشتهر بكونه من أبرز المناهجين عن اللغة العربية والإسلام في القارة الهندية.
14	بعد تلقيه العلم على العديد من المشايخ
15	وتخرج في جامعة ديوبند الإسلامية

REFERENCES

- [1] Abacha, A., & Zweigenbaum, P. (2010). Automatic extraction of semantic relations between medical entities: a rule based approach. International Symposium on Semantic Mining in Biomedicine (SMBM).
- [2] Davis, R., Shrobe, H. & Szolovits, P. 1993. What Is a Knowledge Representation?. AI Magazine, 17-33 .
- [3] Sowa, J. F., & Vivo, M. L. (s.d.). Semantic Networks. Encyclopedia of Artificial Intelligence, Wiley.
- [4] BRACHMANT, R. (1977). What's in a concept: structural foundations for semantic networkst. Int. J. Man-Machine Studies, 127-152.
- [5] Aldhubayi, L., & Alyahya, 2. (2014). Automated Arabic Antonym Extraction Using A Corpus Analysis Tool. Journal of Theoretical and Applied Information Technology.
- [6] Yan , Y., Okazaki, N., Matsuo, Y., & Zhenglu . (2009). Unsupervised Relation Extraction by Mining Wikipedia Texts Using Information from the Web. ACL '09 Proceedings of the Joint Conference of the 47th Annual Meeting of the ACL and the 4th International Joint Conference on Natural Language Processing of the AFNLP, 1021-1029.
- [7] Aronson,, A. (2001). Effective mapping of biomedical text to the UMLS Metathesaurus: the MetaMap program. AMIA Ann Symp Proc.
- [8] Bollegala, D., Matsuo, Y., & Ishizuka, M. (2010,). Relational Duality: Unsupervised Extraction of Semantic Relations between Entities on the Web. International World Wide Web Conference Committee (IW3C2).
- [9] Etzioni, O., Kok, S., Soderland, S., Cafarella, M., Popescu, A., Weld, D., et al. (2004). WebScale Information Extraction in KnowItAll. Proceedings of the 13th international conference on World Wide Web , 100-110.
- [10] Wu, F., & Weld, D. S. (2010). Open Information Extraction usingWikipedia. ACL '10 Proceedings of the 48th Annual Meeting of the Association for Computational Linguistics , 118-127 .
- [11] Lahbib, W., Bounhas, I., & Elayeb, B. (2013). A hybrid approach for arabic semantic relation extraction. International Florida Artificial Intelligence Research Society Conference (FLAIRS 2013).
- [12] الراحي, ع (1998). التطبيق النحوي. الاسكندرية: دار المعرفة الجامعية



ACIT'2017

The International Arab Conference on Information Technology

Yasmine Hammamet, Tunisia

December 22-24, 2017



-
- [13] AlJazeera. . aljazeera.net [Online] Available at: <http://www.aljazeera.net/encyclopedia/icons/2016/6/8/%D9%85%D8%AD%D9%85%D8%AF-%D8%A7%D9%84%D8%A8%D9%86%D9%88%D8%B1%D9%8A-%D8%A7%D9%84%D9%87%D9%86%D8%AF%D9%8A-%D8%A7%D9%84%D9%85%D9%86%D8%A7%D9%81%D8%AD-%D8%B9%D9%86-%D8%A7%D9%84%D8%B6%D8%A7%D8%AF>. [Accessed 28 1 2017].