

Natural Language Generation for Non-Expert Users

Van Duc Nguyen Tran Cao Son Enrico Pontelli

New Mexico State University
Las Cruces, New Mexico, USA

`vnguyen|tson|epontelli@cs.nmsu.edu`

Motivated by the difficulty in presenting computational results—especially when the results are a collection of atoms in a logical language—to users, who are not proficient in computer programming and/or the logical representation of the results, we propose a system for automatic generation of *natural language descriptions* for applications targeting mainstream users. Differently from many earlier systems with the same aim, the proposed system does not employ templates for the generation task. It assumes that there exist some natural language sentences in the application domain and uses this repository for the natural language description. It does not require, however, a large corpus as it is often required in machine learning approaches. The systems consist of two main components. The first one aims at analyzing the sentences and constructs a Grammatical Framework (GF) for given sentences and is implemented using the Stanford parser and an answer set program. The second component is for sentence construction and relies on GF Library. The paper includes two use cases to demonstrate the capability of the system. As the sentence construction is done via GF, the paper includes a use case evaluation showing that the proposed system could also be utilized in addressing a challenge to create an abstract Wikipedia, which is recently discussed in the BlueSky session of the 2018 International Semantic Web Conference.

1 Introduction

Natural language generation (NLG) has been one of the key topics of research in natural language processing, which was highlighted by the huge body of work on NLG surveyed in [21, 7]. With the advances of several devices capable of understanding spoken language and conducting conversation with human (e.g., Google Home, Amazon Echo) and the shrinking gap created by the digital devices, it is not difficult to foresee that the market and application areas of NLG systems will continue to grow, especially in applications whose users are non-experts. In such application, a user often asks for certain information and waits for the answer and a NLG module would return the answer in spoken language instead of text such as in question-answering systems¹ or recommendation systems². The NLG system in these two applications uses templates to generate the answers in natural language for the users. A more advanced NLG system in this direction is described in [16], which works with ontologies annotated using the Attempto language and can generate a natural language description for workflows created by the systems built in the Phylotastic project³. The applications targeted by these systems are significantly different from NLG systems, whose main purpose is to generate high-quality natural language description of objects or reports, such as those reported in the recent AAI conference [12, 6, 15].

The present paper is motivated by the need to generate natural language description of computational results to non-expert users such as those developed in the Phylotastic project. In this project, the users are experts in evolutionary biology but are none experts in ontologies and web services. When a user

¹E.g., the Ergo system: <http://coherentknowledge.com>

²<http://gem.med.yale.edu/ergo/default.htm>

³The Phylotastic project: <http://phylotastic.org>

places a request, he/she will receive a workflow consisting of web services, whose inputs and outputs are specified by instances of classes in the ontologies working with web services, as well as the ordering and relationships between the services. To assist the user in understanding the workflow, a natural language description of the workflow is generated. In order to accomplish the task, the NLG system in the Phylotastic project proposes to annotate elements of the ontologies using Attempto, a simple subset of English with precisely defined syntax and semantics.

In this paper, we propose a system that addresses the limitation of the system discussed in the Phylotastic project [16]. Specifically, we assume that the annotations given in an ontology are natural language sentences. This is a reasonable assumption given that the developers of an ontology are usually those who have intimate knowledge about entities described in the ontology and often have some sort of comments about classes, objects, and instances of the ontology. We then show that the system is very flexible and can be used for the same purpose with new ontologies.

The rest of the paper is organized as follows. Section 2 briefly reviews the basics of Grammatical Framework (GF)[19]. Section 3 describes the main modules of the system. Section 4 includes two use cases of the system using an available ontologies against in the context of reasoning about ontologies. Specifically, it compares with the system used in the Phylotastic project and an ontology about people. This section also contains a use case that highlights the versatility of the proposed system by addressing a challenge to create an abstract Wikipedia [23]. Related works are discussed in Section 5. Section 6 concludes the paper.

2 Background: Grammatical Framework

The Grammatical Framework (GF) [19] is a system used for working with grammars. The GF Resource Grammar Library (RGL)⁴ covering syntax of various languages is the standard library for GF. A GF program has two main parts. The first part is the *Abstract syntax* which defines what meanings can be expressed by a grammar. The abstract syntax defines categories (i.e., types of meaning) and functions (i.e., meaning-building components). An example of an abstract syntax:

Listing 1: Abstract syntax

```
abstract People = {
  flags startcat = Message ;
  cat Message ; People ; Action ; Entity ;
  fun simple_sent : People -> Action -> Entity -> Message ;
    Bill : People; Play : Action; Soccer : Entity ; }
```

Here, Message, People, Action and Entity are types of meanings. startcat flag states that Message is the default start category for parsing and generation. simple_sent is a function accepting 3 parameters, of type People, Action, Entity. This function returns a meaning of Message category. Intuitively, each function in the abstract syntax represents a rule in a grammar. The combination of rules used to construct a meaning type can be seen as a syntax tree.

The second part is composed of *one or more concrete syntax specifications*. Each concrete syntax defines the representation of meanings in each output language. For example, to demonstrate the idea that one meaning can be represented by different concrete syntaxes, we create two concrete syntaxes for two different languages: English and Italian. To translate a sentence to different languages, we only need to provide the strings representing each word in corresponding languages. The GF libraries will take

⁴<https://www.grammaticalframework.org/lib/doc/synopsis/index.html>

responsibility to concatenate the provided strings according to the language grammar to create a complete sentence, which is the representations of the meaning, in the targeted language. The corresponding concrete syntaxes that map functions in the abstract grammar above to strings in English and in Italian is:

Listing 2: Concrete English and Italian syntaxes

```
concrete PeopleEng of People =
open SyntaxEng, ParadigmsEng, ConstructorsEng in {
lincat
  Message = Cl ; People = NP ; Action = V2 ; Entity = NP ;
lin
  simple_sent People Action Entity = mkCl People (mkVP Action Entity) ;
  Bill = mkNP Bill_N; Play = play_V2; Soccer = mkNP soccer_N;
oper
  Bill_N = mkN "Bill" "Bill"; play_V2 = mkV2 "play";
  soccer_N = mkN "soccer"; }

concrete PeopleIta of People =
open SyntaxIta, ParadigmsIta, ConstructorsIta in {
lincat
  Message = Cl ; People = NP ; Action = V2 ; Entity = NP ;
lin
  simple_sent People Action Entity = mkCl People (mkVP Action Entity) ;
  Bill = mkNP Bill_N; Play = play_V2; Soccer = mkNP soccer_N ;
oper
  Bill_N = mkN "Bill"; play_V2 = mkV2 "giocare";
  soccer_N = mkN "calcio" ; }
```

In these concrete syntaxes, the linearization type definition (`lincat`) states that `Message`, `People`, `Action` and `Entity` are type `Cl` (clause), `NP` (noun phrase), `V2` (two-place verb), and `NP` respectively. Linearization definitions (`lin`) indicate what strings are assigned to each of the meanings defined in the abstract syntax. To reduce same string declaration, the operator (`oper`) section defines some placeholders for strings that can be used in linearization. The `mkNP`, `mkN`, `mkV2`, etc. are standard constructors from `ConstructorsEng/Jpn` library which returns an object of the type `NP`, `N` or `V2` respectively.

GF has been used in a variety of applications, such as query-answering systems, voice communication, language learning, text analysis and translation, natural language generation [20, 2], automatic translation⁵.

The translation from English to Italian can be performed as follows in the GF API:

```
parse -lang=PeopleEng "Bill plays soccer" | linearize -lang=PeopleIta
  Bill gioca calcio
```

The above command line produces a syntax tree of the sentence “*Bill plays soccer*” then turn that tree into a `PeopleIta` sentence (in Italian) which is displayed in the second line. Figure 1 shows the meaning in the abstract syntax is represented in Japanese and in Italian, i.e. the two strings represent the same meaning.

⁵ The MOLTO project: <http://www.molto-project.eu>

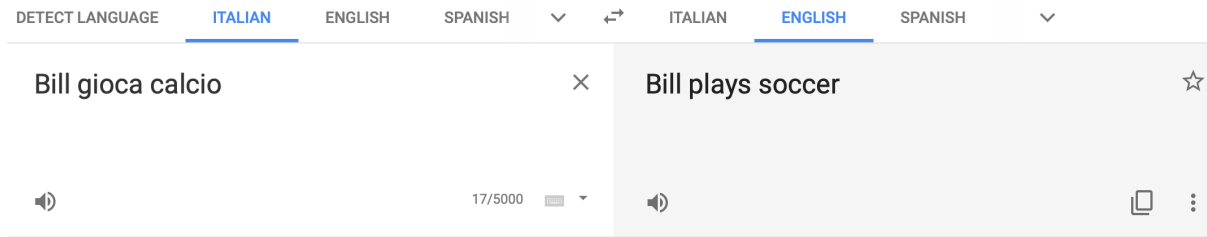


Figure 1: Google translation for the Japanese sentence generated by GF. The two sentences in English and in Italian are the two representations of the meaning encoded in the abstract syntax.

3 Method

To generate a sentence, we need a sentence structure and vocabularies. Our system is developed to emulate the process of a person learning a new language and has to make guesses to understand new sentences from time to time. For example, someone, who understands the sentence “Bill plays a game” would not fully understand the sentence “Bill plays a popular board game” without knowing the meaning of “popular” and “board game” but could infer that the latter sentence indicates that its subject plays a type of game.

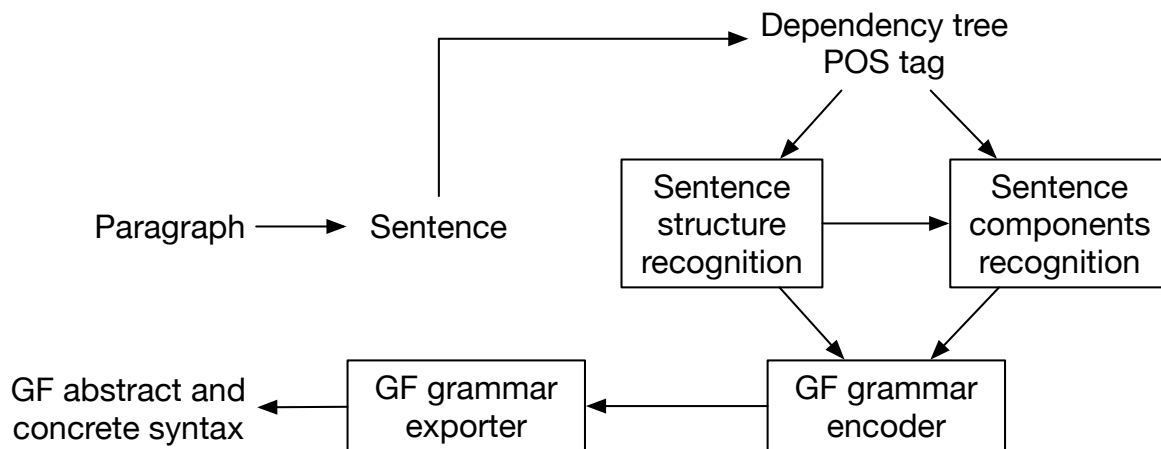


Figure 2: System Overview

The overall design of our system is given in Figure 2. Given a paragraph, our system produces a GF program (a pair of an abstract and a concrete syntax), which can be used for sentence generation. The system consists of two components, understanding sentences and generating GF grammar. The first component is divided into two sub-components, one for recognizing the sentence structure and one for recognizing the sentence components. The second component consists of a GF grammar encoder and a GF grammar exporter. The encoder is responsible for generating a GF grammar for each sentence, while the exporter aggregates the grammars generated from the encoder, and produces a comprehensive grammar for the whole paragraph.

3.1 Sentence Structure Recognition

The sentence structure recognition process involves 2 modules: natural language processing (NLP) module and logical reasoning on result from NLP module. In this paper, we make use of the Stanford Parser tools⁶ described in [4, 17, 22, 14, 10]

The NLP module tokenizes the input free text to produce a dependency-based parse tree⁷ and part-of-speech tag (POS tag). The dependency-based parse tree and the POS tag are then transform into an answer set program (ASP) [8] which contains only facts. Table 1 shows the transformation of the result of NLP module into an ASP program for the sentence “Bill plays a game”. In this table, *nsubj*, *det*, *dobj* and *punct* denote relations in the dependency-based parse tree, and mean *nominal subject*, *determiner*, *direct object* and *punctuation* respectively. Full description of all relations in a dependency-based parse tree can be found in the Universal Dependency website⁸. The second set of notations are the POS tag *PRP*, *VBP*, *DT* and *NN* corresponding to *pronoun*, *verb*, *determiner* and *noun*. Readers can find the full list of POS tag in Penn Treebank Project⁹.

	NLP result	ASP program
Dependency tree	nsubj(plays-2,Bill-1) ROOT(-0,plays-2) det(game-4,a-3) dobj(plays-2,game-4) punct(plays-2,-5)	nsubj(2,1). det(4,3). dobj(2,4). punct(2,5).
POS tag	(Bill, PRP) (plays, VBP) (a, DT) (game, NN) (., .)	pos_tag(1,prp). pos_tag(2,vbp). pos_tag(3,dt). pos_tag(4,nn). pos_tag(5,punct).

Table 1: Transformation from NLP result to asp program

From the collection of the dependency atoms from the dependency-based parse tree, we determine the structure of a sentence using an ASP program, called Π_1 (Listing 3).

Listing 3: Program Π_1

```

structure(1,1) :- nsubj(V,S) .
structure(2,2) :- nsubj(V,S) , dobj(V,O) .
structure(3,3) :- nsubj(V1,S) , xcomp(V1,V2) , dobj(V2,O) .
structure(4,2) :- nsubj(O,S) , cop(O,TOBE) .
structure(5,2) :- nsubjpass(V,S) , auxpass(V,TOBE) .

```

Each of the rule above can be read as *if the right-hand side is true then the left-hand side must be true*. These rules define five possible structures¹⁰ of a sentence represented by the atom *structure(x,y)*. *x* and *y* in the atom *structure(x,y)* denote the type of the structure and the number of dependency relations

⁶<https://nlp.stanford.edu/software/lex-parser.shtml>

⁷https://en.wikipedia.org/wiki/Parse_tree

⁸<http://universaldependencies.org/>

⁹https://www.ling.upenn.edu/courses/Fall_2003/ling001/penn_treebank_pos.html

¹⁰ These are the types of structures that we have implemented in our prototype. Adding additional types will allow us to generate more complicated sentences. This is left for our next work.

applied to activate the rule generating this atom, respectively. We refer to y as the i -value of the structure. For example, $structure(1,1)$ will be recognized if the $nsubj$ relation is in the dependency-based parse tree; $structure(3,3)$ needs 3 dependency relations to be activated: $nsubj$, $xcomp$ and $dobj$. We often use $structure\#x$ to indicate a structure of type x .

Together with the collection of the atoms encoding the relations in the dependency-based parse tree, Π_1 generates several atoms of the form $structure(x,y)$ for a sentence. Among all these atoms, an atom with the highest i -value represents the structure constructed using the highest number of dependency relations. And hence, that structure is the most informative structure that is recognized for the sentence. Observe that $structure(1,1)$ is the most simplified structure of any sentence.

3.2 Sentence Components Recognition

The goal of this step is to identify the relationship between elements of a sentence structure and chunks of words in a sentence from the POS tags and the dependency-based parse tree. For example, the sentence “Bill plays a game” is encoded by a structure #2 and we expect that *Bill*, *plays*, and *game* correspond to the subject, verb, and object, respectively.

We begin with recognizing the main words (components) that play the most important roles in the sentence based on a given sentence structure. This is achieved by program Π_2 (Listing 4). The first four rules of Π_2 determine the main subject and verb of the sentence whose structure is #1, #2, #3, or #5. Structure #4 requires a special treatment since the components following *tobe* can be of different forms. For instance, in “Cathy is gorgeous,” the part after *tobe* is an adjective, but in “Cathy is a beautiful girl,” the part after *tobe* is a noun, though, with adjective *beautiful*. This is done using the four last rules of Π_2 .

Listing 4: Program Π_2

```

2 { sub(S); verb(V) }      :- nsubj(V,S) .
3 { sub(S); obj(O); verb(V) } :- nsubj(V,S), dobj(V,O) .
2 { sub(S); verb(V) }      :- nsubjpass(V,S), auxpass(V,TOBE) .
4 { sub(S); obj(O); verb_1(V1); verb_2(V2) }
   :- nsubj(V1,S), xcomp(V1,V2), dobj(V2,O) .
2 { sub(S); adj(O) }      :- nsubj(O,S), pos_tag(O,jj) .
2 { sub(S); obj(O) }      :- nsubj(O,S), pos_tag(O,nn) .
2 { sub(S); obj(O) }      :- nsubj(O,S), pos_tag(O,nns) .
2 { sub(S); obj(O) }      :- nsubj(O,S), pos_tag(O,cd) .

```

The result of program Π_2 is an one-to-one mapping of some of the words in the sentence into the important components of a sentence, called main components, i.e. subject, object and verb. The mapping is constructed by using the core arguments in Universal Dependency Relations¹¹. Since not every word in the sentence is in a core argument relation, there are some words in the sentence that are not in the domain of the mapping that Π_2 produces. We denote these words are complement components. To identify these words, we encode the Non-core dependents and Nominal dependents from Universal Dependency Relations into the set of rules in program Π_3 .

Program Π_3 (Listing 5), together with the atoms extracted from the dependency-based parse tree such as $compound(P,N)$ (N is compound noun at the position P in the sentence), $amod(P,J)$ (J is an adjective modifier), etc., is used to identify the complement components of the main components computed by Π_2 while maintaining the structure of the sentence created by Π_1 . For example, a complement of a noun could be another noun (as “board” in “board game”), or an adjective (as “popular” in “popular board game”), or a preposition (as “for adults” in “board game for adults”).

¹¹<https://universaldependencies.org/u/dep/>

Structure	GF rules
#1	$NP \rightarrow VP \rightarrow Cl$
#2	$NP \rightarrow V2 \rightarrow NP \rightarrow Cl$
#3	$NP \rightarrow VV \rightarrow V2 \rightarrow NP \rightarrow Cl$
#4	$NP \rightarrow AP \rightarrow Cl$ and $NP \rightarrow NP \rightarrow Cl$
#5	$NP \rightarrow \text{passive}VP \rightarrow Cl$

Table 2: GF Rules Assigned to Each Structure

Listing 5: Program Π_3

```

noun_compound(N)      :- compound(pos, N) .
adj_mod(JJ)           :- amod(pos, JJ) .
noun_conjunction(N)   :- conj(pos, N) .
preposition(COMP, IN) :- nmod(pos, COMP), case(COMP, IN) .
adverbial_modifier(ADV) :- advmod(pos, ADV) .

```

The input of Program Π_3 is the position (*pos*) of the word in the sentence. Program Π_3 is called whenever there is a new complement component discovered. That way of recursive calls is to identify the maximal chunk of the words that support the main components of the sentence. The result of this module is a list of vocabularies for the next steps.

3.3 GF Grammar Encoder

The goal of the encoder is to identify appropriate GF rules for the construction of a GF grammar of a sentence given its structure and its components identified in the previous two modules. This is necessary since a sentence can be encoded in GF by more than one set of rules; for example, the sentence “Bill wants to play a game” can be encoded by the rules

$Bill \rightarrow NP, want \rightarrow VV, play \rightarrow V2, game \rightarrow NP$

and one of the sets of GF rules in the table below:

$V2 \rightarrow NP \rightarrow VP$	$V2 \rightarrow NP \rightarrow VP$
$NP \rightarrow VV \rightarrow VP \rightarrow Cl$	$VV \rightarrow VP \rightarrow VP$
	$NP \rightarrow VP \rightarrow Cl$

In GF, NP, VV, V2, VP, and Cl stand for *noun phrase*, *verb-phrase-complement verb*, *two-place verb*, *verb phrase* and *clause*, respectively. Note that although the set of GF grammatical rules can be used to construct a constituency-based parse tree¹², the reverse direction is not always true. To the best of our knowledge, there exists no algorithm for converting a constituency-based parse tree to a set GF grammar rules. We therefore need to identify the GF rules for each sentence structure.

In our system, a GF rule is assigned to a structure initially (Table 2). Each rule in Table 2 represents the first level of the constituency-based parse tree. It acts as the coordinator for all other succeeding rules.

Given the seed components identified in Section 3.2 and the above GF rules, a GF grammar for each sentence can be constructed. However, this grammar can only be used to generate fairly simple sentences. For example, for the sentence “Bill plays a popular board game with his close friends.”, a GF grammar for structure #2 can be constructed, which can only generate the sentence “Bill plays game.”

¹²Constituency parsing aims to extract a constituency-based parse tree from a sentence that represents its syntactic structure http://nlpprogress.com/english/constituency_parsing.html

For noun components	
$N \rightarrow N \rightarrow CN$	CN: <i>common noun</i>
$N \rightarrow NP$	NP: <i>noun phrase</i>
$AP \rightarrow CN \rightarrow CN$	AP: <i>adjectival phrase</i>
$CN \rightarrow NP$	
$NP \rightarrow Adv \rightarrow NP$	Adv: <i>verb-phrase-modifying adverb</i>
$NP \rightarrow NP \rightarrow ListNP$	
$NP \rightarrow ListNP \rightarrow ListNP$	
$Conj \rightarrow ListNP \rightarrow NP$	Conj: <i>conjunction</i>
For verb components	
$VP \rightarrow Adv \rightarrow VP$	
For adjective components	
$A \rightarrow AP$	A: <i>adjective</i>
$AdA \rightarrow AP \rightarrow AP$	AdA: <i>adjective-modifying adverb</i>

Table 3: Extended GF Rules

because it does not contain any complement components identified in Section 3.2. Therefore, we assign a set of GF rules for the construction of each parameter in the GF rules in Table 2. The set of GF rules has to follow two conventions. The first one is after applying the set of rules to some components of the sentence, the type of the production is one of the type in Table 2, e.g. *NP*, *VP*, *Cl*, *V2*, The second convention is that the GF encoder will select the rules as the order from top to bottom in Table 3. Note that the encoder always has information of what type of input and output for the rule it is looking for.

For instance, we have “game” is the object (main components), and we know that we have to construct “game” in the result GF grammar to be a NP (noun phrase). Program Π_2 identifies that there are two complement components for the word “game”, which are “board” and “popular”, a noun and an adjective respectively. The GF encoder then select the set of rules: $N \rightarrow N \rightarrow CN$ and $A \rightarrow AP$ to create the common noun “board game” and the adjective phrase first. The next rule is $AP \rightarrow CN \rightarrow CN$. The last rule to be applied is $CN \rightarrow NP$. The selection is easily decided since the input and the output of the rules are pre-determined, and there is no ambiguity in the selection process.

The encoder uses the GF rules and the components identified by the previous subsections to produce different constructors for different components of a sentence. A part of the output of the GF encoder for the object “game” is

```
Game = mkNP (mkNP popular_board_game_CN ) (ConstructorsEng.mkAdv
  with_Prep (mkNP close_friend_CN )) ;
```

The encoder will also create the operators that will be included in the `oper` section of the GF grammar for supporting the new constructor. For example, the following operators will be generated for serving the Game constructor above:

```
popular_A = mkA "popular" ;
popular_AP = mkAP popular_A ;
popular_board_game_CN = mkCN popular_AP board_game_N ;
board_game_N = mkN "board game" "board games" ;
close_A = mkA "close" ;
close_AP = mkAP close_A ;
close_friend_CN = mkCN close_AP friend_N ;
friend_N = mkN "friend" "friends" ;
```


3.4 GF Grammar Exporter

The GF Grammar Exporter has the simplest job among all modules in the system. It creates a GF program for a paragraph using the GF grammars created for the sentences of the paragraph. By taking the union of all respective elements of each grammar for each sentence, i.e., categories, functions, linearizations and operators, the Grammar Exporter will group them into the set of categories (respectively, categories, functions, linearizations, operators) of the final grammar.

4 Experiments

We describe our method of generating natural language in two applications. The first application is to generate a natural language description for workflow created by the system built in the Phylotastic project described in [16]. Instead of requiring that the ontologies are annotated using Attempto, we use natural language sentences to annotate the ontologies. To test the feasibility of the approach, we also conduct another use case with the second ontology, that is entirely different from the ontologies used in the Phylotastic project. The ontology¹³ is about people and includes descriptions for certain class.

The second application targets the challenge of creating an abstract Wikipedia from the BlueSky session of 2018 International Semantic Web Conference [23]. We create an intermediate representation that can be used to translate the original article in English to another language. In this use case, we translate the intermediate representation back to English and measure how the translated version stacks up against the original one. We assess the generation quality automatically with BLEU-3 and ROUGE-L (F measure). BLEU [18] and ROUGE [11] algorithms are chosen to evaluate our generator since the central idea of both metrics is “the closer a machine translation is to a professional human translation, the better it is”, thus, they are well-aligned with our use cases’ purpose. In short, the higher BLEU and ROUGE score are, the more similar the hypothesis text and the reference text is. In our use case, the hypothesis for BLEU and ROUGE is the generated English content from the intermediate representation, and the reference text is the original text from Wikipedia.

4.1 NLG for Annotated Ontologies

As described in [16], the author’s system retrieves a set of atoms from an ASP program such as those in Listing 6 where *phylotastic FindScientificNamesFromWeb GET* was shortened to *service*, propagates the atoms, and constructs a set of sentences having similar structure to the sentence “*The input of phylotastic FindScientificNamesFromWeb GET is a web link. Its outputs are a set of species names and a set of scientific names*”. In this sentence, *phylotastic FindScientificNamesFromWeb GET* is the name of the service involved in the workflow of the Phylotastic project. All of the arguments of the atoms above are the names of classes and instances from Phylotastic ontology.

Listing 6: Sample Set of Atoms

```
input(service, web_link).           typeof(web_link, url).
output(service, species_names).     typeof(species_names, names).
output(service, scientific_names).  typeof(scientific_names, names).
```

We replace the original Attempto annotations with the natural language annotations as in Table 4 and test with our system.

¹³Bookmarked URIs in Protege 5.5.0 Build beta-9 or <http://owl.man.ac.uk/2006/07/sssw/people>

Atom	Annotation
input(service, web_link).	The input of <i>service</i> is a web link
output(service, species_names).	The output of <i>service</i> is species names
typeof(web_link, url).	The type of web link is url

Table 4: Atoms from Phylotastic project and its annotation

Tuple	Text
(Kevin,has_pet,Flossie)	Kevin has_pets Flossie.
(Flossie,rdf:type,cow)	Flossie is cow.
(Mick,reads,Daily_Mirror)	Mick reads Daily Mirror.

Table 5: Sample outputs for the people ontology.

With the same set of atoms as in Listing 6, our system generates the following description “*Input of phylotastic FindScientificNamesFromWeb GET is web link. Type of web link is url. Output of phylotastic FindScientificNamesFromWeb GET is scientific names. Output of phylotastic FindScientificNamesFromWeb GET is species names. Type of scientific names is names. Type of species name is names.*”.

We also test our system with the people ontology as noted above. We extract all comments about people and replace compound sentences with simple sentences, e.g., “*Mick is male and drives a white van*” is replaced by the two sentences “*Mick is male*” and “*Mick drives a white van.*” to create a collection of sample sentences. We then use our system to generate a GF program which is used to generate sentences for RDF tuples. Sample outputs for some tuples are in Table 5. This shows that for targeted applications, our system could do a reasonable job.

4.2 Intermediate Representation for Wiki Pages

Since our system creates a GF program for a set of sentences, it could be used as an intermediate representation of a paragraph. This intermediate representation could be used by GF for automatic translation as GF is well-suited for cross-languages translation. On the other hand, we need to assess whether the intermediate representation is meaningful. This use case aims at checking the adequacy of the representation. To do so, we generate the English sentences from the GF program and evaluate the quality of these sentences against the original ones. We randomly select 5 articles from 3 Wikipedia portals: People, Mathematics and Food & Drink.

With the small set of rules introducing in this paper to recognize sentence structure, there would be very limited 4-gram in the generated text appearing in original Wikipedia corpus. Therefore, we use BLEU-3 with equal weight distribution instead of BLEU-4 to assess the generated content. Table 6 shows the summary of the number of assessable sentences from our system. Out of 62 sentences from 3 portals, the system cannot determine the structure 2 sentences in Mathematics due to their complexity. This low number of failure shows that our 5 proposed sentence structures effectively act as a lower bound on sentence recognition module.

In terms of quality, Table 7 shows the average of BLEU and ROUGE score for each portal. Note that the average BLUE score is calculated only on BLEU assessable sentences, while average ROUGE

	People	Mathematics	Food & drink
#sentences	15	24	23
#sentences recognized	15	22	23
BLEU assessable	10	15	11

Table 6: BLEU assessable sentences

	People	Mathematics	Food & drink
BLEU	39.1	33.4	52.2
ROUGE-1	20	17.9	14.6
ROUGE-2	6.7	6.7	5.5
ROUGE-L	11.4	10.5	8.13

Table 7: BLUE and ROUGE score

score is calculated on the sentences whose structure can be recognized and encoded by our system. We note that the BLEU or ROUGE score might not be sufficiently high for a good quality translation. We believe that two reasons contribute to this low score. First, the present system uses fairly simple sentence structures. Second, it does not consider the use of relative clauses to enrich the sentences. This feature will be added to the next version of the system.

Table 8 summarizes the result of this use case. On the left are the paragraphs extracted from the Wikipedia page about Rice¹⁴ in Food & Drink, Decimal¹⁵ in Mathematics, and about Alieu Ebrima Cham Joof¹⁶ from People. As we can see, the main points of the paragraphs are maintained.

5 Related Works

The systems developed in [5, 13, 12] use statistical generation method to produce descriptions of tables or explanation and recommendation from users' reviews of an item. All three systems are capable of generating high quality descriptions and/or explanations. In comparing to these systems, our system does not use the statistical generation method. Instead, we use Grammatical Framework for the generation task. A key difference between these systems and our system lies in the requirement of a large corpus of text in a specific domain for training and generation of these systems. Our system can work with very limited data and a wide range of domains.

Another method for generating natural language explanation for an question-answering system is proposed in [9, 6]. [9] ([9]) describes a system that can give reasonable and supportive evidence to the answer to a question asked to an image, while [6] ([6]) generates explanations for scheduling problem using argumentation. [24] ([24]) use ASP to develop a system answering questions in the do-it-yourself domain. These papers use templates to generate answers. The generated GF program generated by our system, that is used for the NLG task, is automatically created from a provided input.

¹⁴<https://en.wikipedia.org/wiki/Rice>

¹⁵<https://en.wikipedia.org/wiki/Decimal>

¹⁶https://en.wikipedia.org/wiki/Alieu_Ebrima_Cham_Joof

Rice	
Rice is the seed of the grass species <i>Oryza sativa</i> (Asian rice) or <i>Oryza glaberrima</i> (African rice).	Rice is seed of grass species <i>Oryza sativa</i>
As a cereal grain, it is the most widely consumed staple food for a large part of the world's human population, especially in Asia.	it is widely consumed staple food for large part of human population of world in Asia
It is the agricultural commodity with the third-highest worldwide production (rice, 741.5 million tonnes in 2014), after sugarcane (1.9 billion tonnes) and maize (1.0 billion tonnes).	It is agricultural commodity with third-highest worldwide production after sugarcane
Decimal	
The decimal numeral system is the standard system for denoting integer and non-integer numbers.	decimal numeral system is standard system.
It is the extension to non-integer numbers of the Hindu Arabic numeral system.	It is extension to non-integer number of Hindu-Arabic numeral system.
The way of denoting numbers in the decimal system is often referred to as decimal notation.	way is referred to decimal notation.
Alieu Ebrima Cham Joof	
Alieu Ebrima Cham Joof (22 October 1924 – 2 April 2011) commonly known as Cham Joof or Alhaji Cham Joof, (pen name: Alh. A.E. Cham Joof) was a Gambian historian, politician, author, trade unionist, broadcaster, radio programme director, scout master, Pan-Africanist, lecturer, columnist, activist and an African nationalist.	Cham Joof is politician, author, unionist, broadcaster, radio programme director, scout master, Pan-Africanist, lecturer, columnist, activist, African nationalist and Gambian historian.
He advocated for the Gambia's independence during the colonial era.	He advocates for independence of Gambia during colonial era.

Table 8: Original sentences extracted from Wikipedia and corresponding generated sentences

- Educational Advances in Artificial Intelligence, EAAI 2019, Honolulu, Hawaii, USA, January 27 - February 1, 2019*. [1], pp. 2752–2759. Available at <https://aaai.org/ojs/index.php/AAAI/article/view/4126>.
- [7] Albert Gatt & Emiel Krahmer (2018): *Survey of the State of the Art in Natural Language Generation: Core tasks, applications and evaluation*. *J. Artif. Intell. Res.* 61, pp. 65–170, doi:10.1613/jair.5477.
 - [8] Michael Gelfond & Vladimir Lifschitz (1990): *Logic Programs with Classical Negation*. In David H. D. Warren & Péter Szeredi, editors: *Logic Programming, Proceedings of the Seventh International Conference, Jerusalem, Israel, June 18-20, 1990*, MIT Press, pp. 579–597.
 - [9] Shalini Ghosh, Giedrius Burachas, Arijit Ray & Avi Ziskind (2019): *Generating Natural Language Explanations for Visual Question Answering using Scene Graphs and Visual Attention*. CoRR abs/1902.05715. Available at <http://arxiv.org/abs/1902.05715>.
 - [10] Dan Klein & Christopher D. Manning (2002): *Fast Exact Inference with a Factored Model for Natural Language Parsing*. In Suzanna Becker, Sebastian Thrun & Klaus Obermayer, editors: *Advances in Neural Information Processing Systems 15 [Neural Information Processing Systems, NIPS 2002, December 9-14, 2002, Vancouver, British Columbia, Canada]*, MIT Press, pp. 3–10. Available at <http://papers.nips.cc/paper/2325-fast-exact-inference-with-a-factored-model-for-natural-language-parsing>.
 - [11] Chin-Yew Lin & Eduard H. Hovy (2003): *Automatic Evaluation of Summaries Using N-gram Co-occurrence Statistics*. In Marti A. Hearst & Mari Ostendorf, editors: *Human Language Technology Conference of the North American Chapter of the Association for Computational Linguistics, HLT-NAACL 2003, Edmonton, Canada, May 27 - June 1, 2003*, The Association for Computational Linguistics, doi:10.3115/1073445.1073465. Available at <http://aclweb.org/anthology/N/N03/N03-1020.pdf>.
 - [12] Tianyu Liu, Fuli Luo, Qiaolin Xia, Shuming Ma, Baobao Chang & Zhifang Sui (2019): *Hierarchical Encoder with Auxiliary Supervision for Neural Table-to-Text Generation: Learning Better Representation for Tables*. In *The Thirty-Third AAAI Conference on Artificial Intelligence, AAAI 2019, The Thirty-First Innovative Applications of Artificial Intelligence Conference, IAAI 2019, The Ninth AAAI Symposium on Educational Advances in Artificial Intelligence, EAAI 2019, Honolulu, Hawaii, USA, January 27 - February 1, 2019*. [1], pp. 6786–6793. Available at <https://aaai.org/ojs/index.php/AAAI/article/view/4653>.
 - [13] Tianyu Liu, Kexiang Wang, Lei Sha, Baobao Chang & Zhifang Sui (2018): *Table-to-Text Generation by Structure-Aware Seq2seq Learning*. In Sheila A. McIlraith & Kilian Q. Weinberger, editors: *Proceedings of the Thirty-Second AAAI Conference on Artificial Intelligence, (AAAI-18), the 30th innovative Applications of Artificial Intelligence (IAAI-18), and the 8th AAAI Symposium on Educational Advances in Artificial Intelligence (EAAI-18), New Orleans, Louisiana, USA, February 2-7, 2018*, AAAI Press, pp. 4881–4888. Available at <https://www.aaai.org/ocs/index.php/AAAI/AAAI18/paper/view/16599>.
 - [14] Marie-Catherine de Marneffe, Bill MacCartney & Christopher D. Manning (2006): *Generating Typed Dependency Parses from Phrase Structure Parses*. In Nicoletta Calzolari, Khalid Choukri, Aldo Gangemi, Bente Maegaard, Joseph Mariani, Jan Odijk & Daniel Tapias, editors: *Proceedings of the Fifth International Conference on Language Resources and Evaluation, LREC 2006, Genoa, Italy, May 22-28, 2006*, European Language Resources Association (ELRA), pp. 449–454. Available at <http://www.lrec-conf.org/proceedings/lrec2006/pdf/440.pdf.pdf>.
 - [15] Arindam Mitra, Peter Clark, Oyvind Tafjord & Chitta Baral (2019): *Declarative Question Answering over Knowledge Bases Containing Natural Language Text with Answer Set Programming*. In *The Thirty-Third AAAI Conference on Artificial Intelligence, AAAI 2019, The Thirty-First Innovative Applications of Artificial Intelligence Conference, IAAI 2019, The Ninth AAAI Symposium on Educational Advances in Artificial Intelligence, EAAI 2019, Honolulu, Hawaii, USA, January 27 - February 1, 2019*. [1], pp. 3003–3010. Available at <https://aaai.org/ojs/index.php/AAAI/article/view/4157>.
 - [16] Van Nguyen, Tran Cao Son & Enrico Pontelli (2019): *Natural Language Generation from Ontologies*. In: *Practical Aspects of Declarative Languages - 21th International Symposium, PADL 2019, Lisbon, Portugal, January 14-15, 2019, Proceedings*, pp. 64–81, doi:10.1007/978-3-030-05998-9_5.
 - [17] Joakim Nivre, Marie-Catherine de Marneffe, Filip Ginter, Yoav Goldberg, Jan Hajic, Christopher D. Manning, Ryan T. McDonald, Slav Petrov, Sampo Pyysalo, Natalia Silveira, Reut Tsarfaty & Daniel Zeman

- (2016): *Universal Dependencies v1: A Multilingual Treebank Collection*. In Calzolari et al. [3]. Available at <http://www.lrec-conf.org/proceedings/lrec2016/summaries/348.html>.
- [18] Kishore Papineni, Salim Roukos, Todd Ward & Wei-Jing Zhu (2002): *Bleu: a Method for Automatic Evaluation of Machine Translation*. In: *Proceedings of the 40th Annual Meeting of the Association for Computational Linguistics, July 6-12, 2002, Philadelphia, PA, USA.*, ACL, pp. 311–318. Available at <http://www.aclweb.org/anthology/P02-1040.pdf>.
- [19] Aarne Ranta (2004): *Grammatical Framework*. *J. Funct. Program.* 14(2), pp. 145–189, doi:10.1017/S0956796803004738.
- [20] Aarne Ranta (2011): *Grammatical Framework - Programming with Multilingual Grammars*. CSLI Studies in Computational Linguistics, Cambridge University Press. Available at <http://cslipublications.stanford.edu/site/9781575866277.shtml>.
- [21] Ehud Reiter & Robert Dale (1997): *Building applied natural language generation systems*. *Natural Language Engineering* 3(1), pp. 57–87, doi:10.1017/S1351324997001502.
- [22] Sebastian Schuster & Christopher D. Manning (2016): *Enhanced English Universal Dependencies: An Improved Representation for Natural Language Understanding Tasks*. In Calzolari et al. [3]. Available at <http://www.lrec-conf.org/proceedings/lrec2016/summaries/779.html>.
- [23] Denny Vrandečić (2018): *Capturing Meaning: Toward an Abstract Wikipedia*. In Marieke van Erp, Medha Atré, Vanessa López, Kavitha Srinivas & Carolina Fortuna, editors: *Proceedings of the ISWC 2018 Posters & Demonstrations, Industry and Blue Sky Ideas Tracks co-located with 17th International Semantic Web Conference (ISWC 2018), Monterey, USA, October 8th - to - 12th, 2018.*, CEUR Workshop Proceedings 2180, CEUR-WS.org. Available at http://ceur-ws.org/Vol-2180/ISWC_2018_Outrageous_Ideas_paper_6.pdf.
- [24] Yi Wang, Joohyung Lee & Doo Soon Kim (2017): *A Logic Based Approach to Answering Questions about Alternatives in DIY Domains*. In Satinder P. Singh & Shaul Markovitch, editors: *Proceedings of the Thirty-First AAAI Conference on Artificial Intelligence, February 4-9, 2017, San Francisco, California, USA.*, AAAI Press, pp. 4753–4759. Available at <http://aaai.org/ocs/index.php/IAAI/IAAI17/paper/view/14974>.