



An Algorithm for the Reconstruction of Enthymemes for Effective Machine Translation

Enikuomihin Oluwatoyin
Lagos State University, Nigeria
Department of Computer Sci.
Faculty of Science

Odunowo Adebisi T.
Lagos State University, Nigeria
Department of Computer Sci.
Faculty of Science

Mustapha Oluwatoyin S.
Lagos State University, Nigeria
Department of Computer Sci.
Faculty of Science

ABSTRACT

Enthymeme, which is arguments with missing premises or conclusions, is common in natural language text. Enthymeme reconstruction, the art of reformulating arguments with missing propositions, has not been effective in argument classification and consequently, rhetorical Algorithms have yielded poor result. They cannot discover features, text orientation, intent and sentiments in enthymematic arguments. This has led to poor performance of enthymematic Natural language toolkits. Hence, generating new context of enthymematic data reconstruction will provide better and useable insight. The aim of this research is to build a manual annotation framework for enthymemes to enable appropriate tagging and effective classification in argumentation. Manual Annotation technique is used in this experiment to manually separate statements that contain an aspect (enthymemes) from ArguAna corpus of hotel reviews from TripAdvisor.com to know the opinion from the statements. A total of 1201 reviews gave 5575 opinions which were then annotated with defined conclusions. The linear Support Vector Machine (SVM) and fastText classifier were used to train and test data while Valence Aware Dictionary for sEntiment Reasoning (VADER) was used to assign scores for each word based on sentiments. MATLAB and Python programming language were used for model implementation. The supervised learning approach showed the best performance results on the test set with a macro averaged F1-scores of 0.72 and 0.94 for explicit and implicit stances respectively. The identified implicit stances are explicit premises of either complete arguments or enthymemes. (If they are premises of complete arguments, there are other, additional premises.) The identified explicit stances can represent common knowledge information for the implicit premises, thus becoming explicit premises to fill in the gap present in the respective enthymemes. The experimental framework shows that manual annotation of enthymeme data can provide better and useable insight in machine based annotation.

General Terms

Reconstruction of Enthymemes and Machine Translation

Keywords

Arguments, Enthymemes, Manual Annotation, Machine Translation, Rhetorical Algorithm, Syllogism

1. INTRODUCTION

Argumentation which is the process by which arguments are constructed and handled has four main tasks to undertake: identification, analysis, evaluation and invention. Identification is the task of determining the conclusion, premises and scheme of an argument from natural discourse. Argumentation comes from the question “Was the right decision made? Was it well founded?” For every decision

made, one might be asked to justify, explain or defend how it was arrived at [1]. An argument consists of two or more propositions, one proposition functions as the claim (also known as the conclusion), and a set of one or more propositions serve as supports (also known as premises). An argument has a structure, which plays a key role in determining the presence or absence of an argument [2]. Argumentation theory is an area of study including formal, semi-formal, and informal methods for the identification, analysis, evaluation, and production of human arguments, methods which generally go beyond formal logic [3]. The conclusion of one argument may take on the role of premise in another argument or a premise of one argument may also be a premise of another argument, so that arguments about a particular topic are interconnected and their interactions may be viewed as a graph, with nodes representing argument components and directed arcs representing relationships between them, such as forms of inferences. Real arguments are often enthymemes instead of completely specified deductive arguments. This means that some parts of the pair (support, claim) may be missing because they are supposed to belong to some “common knowledge”, and then should be deduced by the agent which receives the enthymeme[4].

In logic, an enthymeme is said to be an argument, or chain of argumentation, with one or more missing (implicit) premises or conclusions [5]. One of the most effective and oldest tools available to a rhetorician is the enthymeme. Enthymeme is a rhetorical algorithm, an incomplete syllogism whose implicit, unstated completion is realized only when an audience is persuaded to perform that completion—optimally in a manner aligned with a rhetorician’s intent [6]. These arguments with missing (unstated) premise or conclusions are reconstructed using deductive logic (like syllogism). Words with syllogism have both major (general statement) and minor (more specific statement) premises to form the conclusion. For example, “**All birds lay eggs. A swan is a bird. Therefore, a swan lays eggs**”. In this example, the major premise is that all birds lay eggs. The minor premise is that a swan is a bird. The inference relates these two premises to conclude that if a swan is a bird it must lay eggs. In an enthymeme, one of the premises—major or minor—is implied and thus left out of the reasoning. Even the conclusion can be omitted in an enthymeme because it is obvious enough to the reader or listener. For example, **We are dependent; therefore we should be humble. The complete syllogism would be: Dependent creatures should be humble; we are dependent creatures; therefore we should be humble.** Decisions may be connected with the construction of rules which are informative to the subject. One way of achieving such a representation is to find correlations between particular sub-symbolic features and natural predicates which permit us to illustrate conditions which align with the network’s output. When we want to provide the required information, we are



indeed, producing an enthymeme [7]. We can find arguments almost everywhere: scientific texts, legal texts and court decisions, TV adverts, biomedical texts, patents, reviews, debates, dialogs, news, and so on [8] hence, adopting Computational approaches such as the use of enthymemes in text summarization; question answering system, autonomous machines, machine translations, refinement of arguments, discourse analysis and legal support systems etc. for achieving reliable results to reconstruct the enthymemes in an argument for effective machine translation will be the focus of this paper since Enthymeme reconstruction, the art of reformulating arguments with missing propositions, has not been effective in text summarization, question answering system, autonomous machines, machine translations, refinement of arguments, information retrieval processes, discourse analysis, children news rendering, counterfactual sentences and legal support systems etc. and consequently, rhetorical Algorithms have yielded poor result. They cannot discover features, text orientation, intent and sentiments in enthymematic arguments. This has led to poor performance of enthymematic Natural Language Toolkits. Generating new context of manual annotation of enthymeme data selection and reconstruction to provide better and useable insight in machine-based annotation for reconstruction of Enthymemes is our goal. This paper focuses on developing a Manual Annotation technique to manually separate statements that contain an aspect (enthymemes) from ArguAna corpus to know the opinion from the statements [9]. The linear SVM (Support Vector Machine) classifier is engaged to train and test data for effectiveness [10] of arguments structures whether explicit or implicit using the n-grams and Part of Speech (POS) tags. Also, VADER (Valence Aware Dictionary for sEntiment Reasoning) lexical resource will be used to assign scores for each word based on sentiments [11].

2. RELATED WORKS

Argumentation has become an Artificial Intelligence keyword for years, especially for handling inconsistency in knowledge bases, for decision making under uncertainty, and for modeling interactions between agents [12]. Recent work on argument interpretation and enthymeme reconstruction includes that of [13] which classify argumentation schemes using explicit premises and conclusion on the Araucaria dataset as a proposal to reconstruct enthymemes. The argumentation scheme classification system presented in the work paper introduces a new task in research on argumentation. According to Lippi and Torroni (2015a) one particularly important aspect of argumentation mining is claim identification. Most of the current approaches are engineered to address specific domains. However, argumentative sentences are often characterized by common rhetorical structures, independently of the domain [1]. The approach thus propose (Context-Independent Claim Detection for Argument Mining) a method that exploits structured parsing information to detect claims without resorting to contextual information, and yet achieve a performance comparable to that of state-of-the-art methods that heavily rely on the context. [14] defined the challenging task of automatic claim detection in a given context and discuss its associated unique difficulties using Context Dependent Claim Detection (CDCD). Recently, efforts in arguments mining has focused on extracting arguments pertaining to a specific domain such as online debates. [15] made a step towards argument-based opinion mining from online discussions (user comments on blogs and forums) and introduce a new task of

argument recognition. Matching user-created comments to a set of predefined topic-based arguments, which can be either attacked or supported in the comment. They present a manually-annotated corpus for argument recognition in online discussions. [16] This report summarizes the objectives and evaluation of the SemEval 2015 task on the sentiment analysis of figurative language on Twitter (Task 11). This is the first sentiment analysis task wholly dedicated to analyzing figurative language on Twitter. [17] found that the ability to analyze the adequacy of supporting information is necessary for determining the strength of an argument. This is especially the case for online user comments, which often consist of arguments lacking proper substantiation and reasoning. Thus, they develop a framework for automatically classifying each proposition as UNVERIFIABLE, VERIFIABLE NONEXPERIENTIAL, or VERIFIABLE EXPERIENTIAL, where the appropriate type of support is reason, evidence, and optional evidence, respectively. Our efforts had been centered on using the stance and the context of the relevant opinion to help in detecting and reconstructing enthymemes present in a specific domain of online reviews. Lippi and Torroni (2015b) address the domain-dependency of previous work by identifying claims that are domain-independent by focusing on rhetoric structures and not on the contextual information present in the claim. [18] argue that an annotation scheme for argumentation mining is a function of the task requirements and the corpus properties. There is no one-size fits-all argumentation theory to be applied to realistic data on the Web. To formalize the problem of annotating arguments, they apply a Claim-Premise scheme, and in the second case, modified Toulmin's scheme. The finding reveals that the choice of the argument components to be annotated strongly depends on the register, the length of the document, and inherently on the literary devices and structures used for expressing argumentation. [19] reports that If it can be detected whether a premise is missing in an argument, then we can either fill the missing premise from similar/related arguments, or discard such enthymemes altogether and focus on complete arguments by drawing a connection between explicit vs. implicit opinion classification in reviews, and detecting arguments from enthymemes. Effort at identifying enthymemes includes that of Feng and Hirst (2011) which classify argumentation schemes using explicit premises and conclusion on the Araucaria dataset, which they propose to use to reconstruct enthymemes. Similar to (2011), Walton (2010) investigated how argumentation schemes can help in addressing enthymemes present in health product advertisements. Amgoud et al., (2010) propose a formal language approach as an extension of Dung's abstract argumentation system to construct arguments from natural language texts that are mostly enthymemes. Their work is related to mined arguments from texts that can be represented using a logical language and our work could be useful for evaluating [12] on a real dataset. Our contribution to the field of argumentation is to develop a manual annotation model that classifies stances which can identify enthymemes and implicit premises that are present in enthymematic documents.

3. METHODOLOGY

3.1 Dataset

The dataset for this research are English Language corpus acquired from ArguAna corpus [20] of hotel service reviews from TripAdvisor.com to manually separate enthymematic statements i.e. statements that contained an aspect (enthymemes) from those that did not contain an aspect. Most

statements/arguments used as dataset though simple and can be reconstructed with ease, directly refers to certain aspects of the hotel or directly to it, the rest were discarded because it will require much deeper analysis to construct such arguments that will be relevant to the hotel. Each argument in the data set majorly has statements that are either positive/negative statement serving as the claim, and others serving as premises. Hence, Manual separation was made possible because the information needed was available in the corpus, though this can as well be achieved automatically using opinion mining tools but will not be as precise and accurate as manual. Obtained were 1201 annotated statements which gave 5575 opinions with the four categories (Issue, Blame, Appreciate, Call for Action) on the speeches. A speech can be labeled with multiple categories as members can appreciate and raise issues in the same speech. Arguments used were classified as positive, negative and neutral. Statements with a positive or negative sentiment were more opinion oriented and hence discarded were the statements that annotated as neutral. The definitions and examples of the four categories are explained in the below table respectively.

Table 1: Definition of the examples

Examples	Count
Issue	<i>Raise problems in general which need attention.</i>
Blame	<i>Blaming the hotel owner or the managers of the hotel.</i>
Appreciate	<i>Appreciating and justifying good hotels and the benefits they got.</i>
Call for Action	<i>Speeches in which members suggest, request for new infrastructures.</i>

The reviews were based on possible (predefined) conclusions for the hotel reviews which were either:

Conclusion 1: In favour of an aspect of the hotel or the hotel itself.

Conclusion 2: Against an aspect of the hotel or the hotel itself.

We then annotated each of the 5575 opinions with one of these conclusions to make the annotation procedure easier, since each opinion related to the conclusion forms either a complete argument or an enthymeme. During the annotation process, each opinion was annotated as either explicit or implicit based on the stance definitions given above.

3.2 Manual Annotation Process

Manual annotation is a well-known framework in Natural Language Processing (NLP). It involves adding labels of linguistic nature or reflecting the usage of NLP technologies on some oral or written discourse [21]. There are two basic processes for adding data--about--data: automatic and manual. Automatic annotation is less precise but can operate over many more documents than humans can reasonably address. Manual is more precise (to a point), but very labor--intensive and is often use to train a machine to perform automatic annotation [22]. Manual annotations vary in nature (phonetic, morpho-syntactic, semantic or task-oriented labels), in the range they cover (they can concern a couple of characters, a word, a paragraph or a whole text), in their

degree of coverage (all the text is annotated or only a part of it) and in their form (atomic value, complex feature structures or relations and even cross-document alignment relations).

3.3 Explicit/Implicit Opinions and Arguments/Enthymemes

A firm relationship exists between detecting whether a particular statement carries an explicit or an implicit opinion, and whether there is a premise that supports the conclusion (resulting in an argument) or not (resulting in an enthymeme). For example:

the following two statements A1 and A2:

A1 = I am awfully disappointed with the room.

A2 = The room is small.

Both statements above express a negative sentiment towards the room aspect of this hotel. In A1, the position of the reviewer (whether the reviewer is in favour or against the hotel) is explicitly stated by the phrase extremely disappointed. Consequently, we refer to A1 as an explicitly opinionated statement about the room. However, to interpret A2 as a negative judgment we must possess the knowledge that being small is often considered as negative with respect to hotel rooms, whereas being small could be positive with respect to some other entity such as a mobile phone. The stance of the reviewer is only implicitly conveyed in A2. A2 can be referred as an implicitly opinionated argument about the room. Given the conclusion that this reviewer did not like this room (possibly explicitly indicated by a low rating given to the hotel), the explicitly opinionated statement A1 would provide a premise forming an argument, whereas the implicitly opinionated statement A2 would only form an enthymeme. Thus:

Argument

Major premise: I am awfully disappointed with the room.

Conclusion: The reviewer is not in favour of the hotel.

whereas:

Enthymeme

Major premise: A small room is considered bad (unstated).

Minor premise: The room is small.

Conclusion: The reviewer is not in favour of the hotel.

3.3.1 Algorithm

The framework for enthymeme detection via opinion classification is expressed below in two major steps. This assumes a separate process to extract the ("predefined") conclusion, for example from the rating that the hotel is given.

Step-1: Opinion structure extraction

- Extract statements that express opinions with the help of local sentiment (positive or negative) and discard the neutral statements.
- Perform an aspect-level analysis to obtain the aspects present in each statement and those statements that include an aspect are considered and the rest of the statements (neutrals) are discarded.
- Classify the stance of statements as being explicit or implicit.

Step-2: Premise extraction

- Explicit opinions paired with the predefined

conclusions can give us complete arguments.

- b. Implicit opinions paired with the predefined conclusions can either become arguments or enthymemes. Enthymemes require additional premises to complete the argument.
- c. Common knowledge can then be used to complete the argument.

3.4 Text Classification

To differentiate whether the stances are explicit or implicit opinions, the work was classified as a binary problem with the following features to determine the sentiment in each statement reviewed.

Baseline: As a baseline comparison, statements containing words from a list of selected cues such as excellent, great, worst etc. are predicted as explicit and those that do not contain words present in the cue list are predicted as implicit. The criteria followed is that the statement should contain atleast one cue word to be predicted as explicit.

N-grams (Uni, Bi): Unigrams (each word) and bigrams (successive pair of words).

Part of Speech (POS): The Natural Language Tool Kit tagger helps in tagging each word with its respective part of speech tag and we use the most common tags (noun, verb and adjective) present in the explicit opinions as features.

Part of Speech (POS Bi) As for POS, but we consider the adjacent pairs of part of speech tags as a feature.

VADER: (Valence Aware Dictionary for sEntiment Reasoning) [11] which is a model used for automatic text sentiment analysis that is sensitive to both polarity (positive/negative) and intensity (strength) of emotion was used to detect the polarity of each speech. The tool uses a simple rule-based model for general sentiment analysis and generalizes more favorably across contexts than any of many benchmarks such as LIWC and SentiWordNet. The tool takes the input as a sentence and gives a score between -1 and 1. The polarity of a speech is calculated by taking the sum of the polarities of the sentences. If the sum is greater than zero, then it is classified as positive, if it is less than zero, then it is classified as negative and if it is equal to zero then it is classified as neutral.

Sentiment score is measured on a scale from -4 to +4, where -4 is the most negative and +4 is the most positive. The midpoint 0 represents a neutral sentiment. Individual words have a sentiment score between -4 to 4, but the returned sentiment score of a sentence is between -1 to 1. However, we apply a normalization to the total to map it to a value between -1 to 1.

The normalization used by Hutto is

$$\frac{x}{\sqrt{x^2 + \alpha}}$$

where^x is the sum of the sentiment scores of the constituent words of the sentence and is a normalization parameter that was set to 15.

3.5 Classification Modelling

SVM and fastText which is a text library in SVM was used for both the classification tasks and preliminary experiments. The text was pre-processed by removing the punctuation and

lowering the case. The reason for using fastText is because of its promising results using n-grams features. The detection of the categories that was developed also largely dependend on the lexicons and so fastText and SVM with word related features is better way to go for preliminary experiments. For example, the category appreciate can be easily be identified using lexicons such as congratulate, appreciate, benefit etc. The training and testing data was divided in the ratio of 80:20 for classification. Accuracy was recorded for each class as it is a multi-label classification problem.

Table 2: Accuracy Score for Classification Task

Task /Metric	fastText	SVM
Call for Action	0.745	0.72
Issue	0.604	0.56
Blame	0.783	0.84
Appreciate	0.679	0.62

4. EXPERIMENTS AND RESULT

Text classification is the core task to many applications, hate speech detection, emotion detection, and sentiment analysis and in dialogue systems. After thorough investigation of many speeches it was found that the statements made by reviewers can be used for detecting few of the above mentioned classification tasks.

4.1 Annotation

As a preliminary step, four major categories of the statements by the reviewers were created. The definitions and examples of the four categories are explained in the below table. Below are quote portions of a few speeches which will give an idea of the data being presented:

“Having booked a deluxe room at the Pulitzer, we were very much looking forward to our one night stay. However, from the moment that we had to walk through a series of corridors that got darker and more tatty the further we went, our expectations were disappointed”

Table 3: Example statements of the categories

Categories	Count
Issue Issue	Having lugged our suitcases up a winding and dirty staircase we got to our room. It was small and smelt very strongly of smoke. When booking, I had ticked the non-smoking preference and neither of us felt comfortable in such a smoky atmosphere.....
Blame Blame	The policy of the hotels management is going in one direction and the customers which come in are not happy with the services.....
Appreciate Appreciate	The reviews are correct, an excellent place to stay The Nob Hill Motor Inn, just fantastic. If you're going to San Francisco just stays there. Possibly the best value hotel in the world, free parking in San Francisco only two blocks from a cable car. This sort of hotel is rare....

Cal Call for Action	We had made a very early fully refundable reservation to stay at the Hotel Pulitzer through a travel agent but had the travel agent cancel the reservation a couple of months before our arrival date because we had a found a better deal. Though our travel agent promptly cancelled the reservation, the Pulitzer still charged us and refused to credit us despite the fact that they could provide no documentation to explain their position.....
---------------------------	---

The annotator agreement is shown in Table 4 and is evaluated using two metrics, one is the Kohen's Kappa and other is the inter-annotator agreement which is the percentage of overlapping choices between the annotators.

Table 4: Inter Annotator agreement metrics of Annotated Data

Categories	Kohen's Kappa	Inter Annotator Agreement
Issue	0.67	0.84
Blame	0.65	0.90
Appreciate	0.88	0.94
Call for Action	0.64	0.92

Table 5: Interpretation of Fleiss' Kappa Scores for Annotator Agreement.

Fleiss' Kappa	Interpretation
≤ 0	Poor agreement
0.01 - 0.20	Slight agreement
0.21 - 0.40	Fair agreement
0.41 - 0.60	Moderate agreement
0.61 - 0.80	Substantial agreement
0.81 - 1.00	Almost perfect agreement

The inter annotator agreement for the stance categories were 0.92. The high values of inter annotator scores clearly explain how easy it was to delineate each category and annotate them. It also signifies that the definition of the category that needed to be annotated, were very clear. Also, according to the Table 5, most of the agreements come under substantial agreements.

4.2 Keywords and Summarization

To enrich the dataset, automatically generated summary was added as a key value pair along with the speeches of the debate to enrich the dataset. TextRank which is an extractive summarizer for summarizing the entire debate and for finding keywords in the debate have been used. TextRank is a graph based ranking model for text processing specifically keywords Extraction and Sentence Extraction. TextRank performs better in text summarization using graph based techniques. Added were these two extra fields i.e. the keywords extracted by TextRank and the summary created by TextRank in the debates collection.

4.2.1 Detection of Polarity

To detect the polarity of each speech, VADER [11] an automatic sentiment analysis tool was used. The tool uses a simple rule-based model for general sentiment analysis and generalizes more favorably across contexts than any of many benchmarks such as LIWC and SentiWordNet. The tool takes the input as a sentence and gives a score between -1 and 1. The polarity of a speech is calculated by taking the sum of the polarities of the sentences. If the sum is greater than zero, then it is classified as positive, if it is less than zero, then it is classified as negative and if it is equal to zero then it is classified as neutral. The statistics of the data is presented in Table 6.

Table 6: Sentiment Polarity of Speeches

Categories	Count
Positive	4006
Negative	1457
Neutral	112
Total	5575

4.3 Stance Classification

In this section, two tasks was dealt with, task one is the classification of the stances the speakers have taken and task two is the detection of few classes such as blame, call for action, appreciate and issue as said earlier. Stance classification differs from sentiment analysis. For instance, the number of speeches that were annotated as for i.e. 919 had only 719 labelled as positive and the number of speeches that were annotated as against i.e. 282 had only 89 as negatively labelled. So, these statistics clearly indicate the difference between polarity detection and stance classification. fastText and SVM was used for preliminary experiments. Pre-processing of the text by removing the punctuation and lowering the case was done. The reason using fastText is because of its promising results uses n-grams features. The detection of categories that had been developed also largely dependent on the lexicons and so fastText and SVM with word related features is better way to go for preliminary experiments. For example, the category appreciate can be easily be identified using lexicons such as congratulate, appreciate, benefit etc. The training and testing data was divided in the ratio of 80:20 for classification. As mentioned above we used fastText and SVM for both the classification tasks. Accuracy was recorded for each class as it is a multi-label classification problem. The results are shown in Table 7 and Table 8. Also, the parameters used for fastText is described in Table 9

Table 7: Accuracy Score for Classification Task 1

Task/Metric	fastText	SVM
For/against	0.80	0.76

Table 8: Accuracy Score for Classification Task 2

Task 2/Metric	fastText	SVM
Call for Action	0.745	0.72
Issue	0.604	0.56
Blame	0.783	0.84
Appreciate	0.679	0.62

Table 9: Precision, Recall & F1 scores of fastText for Classification Task 2

Task 2/Metric	Precision	Recall	F1 score
Call for Action	0.76	0.95	0.85
Issue	0.51	0.74	0.61
Blame	1.0	0.78	0.87
Appreciate	0.80	0.69	0.74

Linked Argument

The entities are related using the support and attack relations. A premise has an outgoing directional relation towards a claim or a premise (in case of a serial argument) and a claim only has an incoming directional relation from the Premise. The below diagram explains annotation schema with a brief explanation of the claim and premises from the dataset.

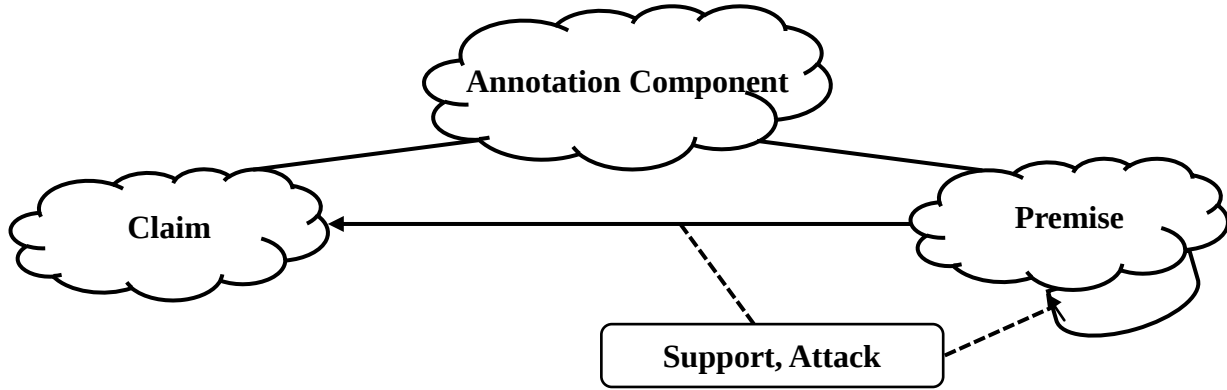


Figure 1: Annotation Schema

Annotation of claims achieves higher agreement values than that of premises in this dataset. The boundary agreement between the annotators is also higher for claims than that of premises according to the values obtained. The reason is that premises are a bit difficult to identify as they may span anywhere across the speech as discussed in the mistakes section above.

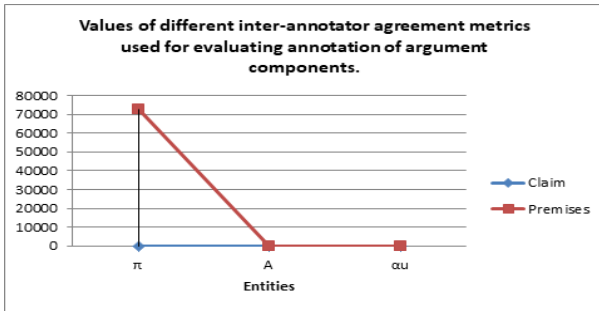


Figure 2: Values of different inter-annotator agreement metrics used for evaluating annotation of argument components

4.4 Classification Modeling

Support Vector Machine (SVM)

SVM classifier is a function (Equation 1) which at a conceptual level is often called the Optimal Margin Classifier. It converts the data into a linearly separable model such that the review outcome is done depending on where the data is placed on a hyperplane

$$h_{w,b}(x) = g(w^T x + b) \dots \text{Equation 1}$$

Where x is the number of features in our dataset, w is weight of the hyperplane. This is a regression equation in a straight-line which linearly divides our dataset into positive and negative. That is, class labels are denoted as -1 (non-defective) for negative class and $+1$ for positive class (defective), i.e. $y \in$

$\{-1, +1\}$. The functional and geometric margin SVM hyperplane is given as: $w^T x^{(i)} + b = 0$. This is the equation of separating the hyperplane. For the training set, the functional margin is given as $\hat{y}^{(i)} = y^{(i)}(x^{(i)} + b)$.

The model aims at maximizing the geometric margin and returning the corresponding hyperplane. The dataset was collected because it contains measurable features in enthymeme detection in natural language texts processes. By preprocessed the dataset and partitioned into training and testing using a 3-fold cross-validation. SVM model was applied and the accuracy, specificity and sensitivity of our model were computed.

5. CONCLUSION

The effort in this paper is geared towards the manual annotation approach to the reconstruction of enthymemes for effective machine translation features in natural language tasks and testing it on the summarization task. A manual annotation approach was designed inspired from previous research techniques on reconstruction of enthymemes. The stages involved in the approach were duly explained and evaluated using SVM metrics with other high performing extractive summarization techniques. The focus of this work is on a specific domain of online reviews and propose an approach that can help in enthymemes detection and reconstruction. Recall, online reviews contain aspect-based statements that can be considered as stances representing for/against views of the reviewer about the aspects present in the product or service and the product/service itself. The proposed approach is a two-step approach that detects the type of stances based on the contextual features, which can then be converted into explicit premises, and these premises with missing information represents enthymemes. Also proposed is a solution using the available data to represent common knowledge that can fill in the missing information to complete the arguments. The first-step requires automatic detection of the stance types— explicit and implicit, which had been covered in this paper. The supervised learning



approach is use to classify the stances using a SVM classifier, the best performance results on the test set with a macro averaged F1-scores of 0.72 and 0.94 for explicit and implicit stances respectively. Here, the identified implicit stances are can serve as explicit premises of either complete arguments or enthymemes. (If they are premises of complete arguments, there are other, additional premises.) The identified explicit stances can then represent common knowledge information for the implicit premises, thus becoming explicit premises to fill in the gap present in the respective enthymemes.

6. ACKNOWLEDGMENTS

Our thanks goes to all authors and intellects whose write-ups and contribution to knowledge had been considered in the making of this paper.

7. REFERENCES

- [1] Lippi, M., and Torroni, P.: 'Context-independent claim detection for argument mining', Twenty-Fourth International Joint Conference on Artificial Intelligence, 2015
- [2] Sergeant, A.: 'Automatic argumentation extraction', in Editor (Ed.)^(Eds.): 'Book Automatic argumentation extraction' (Springer, 2013, edn.), pp. 656-660
- [3] Mebane, W.: 'Detection of claims and supporting evidence in wikipedia articles on controversial topics', 2017
- [4] Mailly, J.-G.: 'Using enthymemes to fill the gap between logical argumentation and revision of abstract argumentation frameworks', arXiv preprint arXiv:1603.08789, 2016
- [5] Walton, D.: 'The three bases for the enthymeme: A dialogical theory', *Journal of Applied Logic*, 2007, 6, (2008), pp. 361-379
- [6] Brock, K.: 'Enthymeme as Rhetorical Algorithm', *Present Tense Journal*, 2014, 4, (1), pp. 2 - 8
- [7] Pedersen, T., and Dyrkolbotn, S.K.: 'The legally mandated approximate language about AI', *Open Journal Systems*, 2018
- [8] da Rocha, G.F.: 'ArgMine: Argumentation Mining from Text', 2016
- [9] Hu, M., and Liu, B.: 'Mining opinion features in customer reviews', in Editor (Ed.)^(Eds.): 'Book Mining opinion features in customer reviews' (2004, edn.), pp. 755-760
- [10] Wachsmuth, H., Trenkmann, M., Stein, B., and Engels, G.: 'Modeling review argumentation for robust sentiment analysis', in Editor (Ed.)^(Eds.): 'Book Modeling review argumentation for robust sentiment analysis' (2014, edn.), pp. 553-564
- [11] Hutto, C.J., and Gilbert, E.: 'Vader: A parsimonious rule-based model for sentiment analysis of social media text', in Editor (Ed.)^(Eds.): 'Book Vader: A parsimonious rule-based model for sentiment analysis of social media text' (2014, edn.), pp.
- [12] Amgoud, L., and Besnard, P.: 'A formal analysis of logic-based argumentation systems', in Editor (Ed.)^(Eds.): 'Book A formal analysis of logic-based argumentation systems' (Springer, 2010, edn.), pp. 42-55
- [13] Feng, V.W., and Hirst, G.: 'Classifying arguments by scheme', in Editor (Ed.)^(Eds.): 'Book Classifying arguments by scheme' (Association for Computational Linguistics, 2011, edn.), pp. 987-996
- [14] Levy, R., Bilu, Y., Hershovich, D., Aharoni, E., and Slonim, N.: 'Context dependent claim detection', in Editor (Ed.)^(Eds.): 'Book Context dependent claim detection' (2014, edn.), pp. 1489-1500
- [15] Boltužić, F., and Šnajder, J.: 'Back up your stance: Recognizing arguments in online discussions', in Editor (Ed.)^(Eds.): 'Book Back up your stance: Recognizing arguments in online discussions' (2014, edn.), pp. 49-58
- [16] Ghosh, A., Li, G., Veale, T., Rosso, P., Shutova, E., Barnden, J., and Reyes, A.: 'Semeval-2015 task 11: Sentiment analysis of figurative language in twitter', in Editor (Ed.)^(Eds.): 'Book Semeval-2015 task 11: Sentiment analysis of figurative language in twitter' (2015, edn.), pp. 470-478
- [17] Park, J., and Cardie, C.: 'Identifying appropriate support for propositions in online user comments', in Editor (Ed.)^(Eds.): 'Book Identifying appropriate support for propositions in online user comments' (2014, edn.), pp. 29-38
- [18] Habernal, I., Eckle-Kohler, J., and Gurevych, I.: 'Argumentation Mining on the Web from Information Seeking Perspective', in Editor (Ed.)^(Eds.): 'Book Argumentation Mining on the Web from Information Seeking Perspective' (2014, edn.), pp.
- [19] Rajendran, P., Bollegala, D., and Parsons, S.: 'Contextual stance classification of opinions: A step towards enthymeme reconstruction in online reviews', in Editor (Ed.)^(Eds.): 'Book Contextual stance classification of opinions: A step towards enthymeme reconstruction in online reviews' (2016, edn.), pp. 31-39
- [20] Wachsmuth, H., Trenkmann, M., Stein, B., Engels, G., and Palakarska, T.: 'A review corpus for argumentation analysis', in Editor (Ed.)^(Eds.): 'Book A review corpus for argumentation analysis' (Springer, 2014, edn.), pp. 115-127
- [21] Fort, K., Nazarenko, A., and Rosset, S.: 'Modeling the complexity of manual annotation tasks: a grid of analysis', in Editor (Ed.)^(Eds.): 'Book Modeling the complexity of manual annotation tasks: a grid of analysis' (2012, edn.), pp.
- [22] Petrillo, M., and Baycroft, J.: 'Introduction to manual annotation', Fairview Research, 2010