

You Said it? How Mis- and Disinformation Tweets Surrounding the Corona-5G-conspiracy Communicate Through Implying

Lynn de Rijk, Radboud University, Netherlands

l.derijk@student.ru.nl

ABSTRACT

This paper aims to investigate if implied meaning plays a role in mis-/disinformation tweets and what linguistic cues might signal this. A qualitative analysis of 130 mis-/disinformation tweets regarding the corona-5G-conspiracy using Speech Act Theory, shows that often meaning is implied by leaving out coherence markers, putting the words in someone else's mouth through citing and ambiguous phrasing/punctuation.

1 INTRODUCTION

Systems that automatically recognize mis-/disinformation are challenged by certain basic communicative features, such as sarcasm or implied meaning. It is therefore relevant to find out if indirect communication plays a role in mis-/disinformation and how large this role is. Speech Act Theory [1] can be used as a framework to index implied meaning. It distinguishes three forces in every utterance: 1) locution, what is said literally, 2) illocution, what the utterance does and 3) perlocution, what happens as a result. For example, in: "Peter, you are standing on my foot", the locutionary force is asserting this state of affairs. The illocution, would generally be requesting that this Peter lifts his foot, now that he is made aware. The perlocutionary act would then be that Peter indeed places his foot elsewhere.

However, Micheal Geis argues that that the illocutionary force of a Speech Act (SA) is mostly dependent on the context, not linguistic cues [2]. Geis gives the example of a teacher asking a student if they can solve a quadratic equation vs. the same question being asked by a fellow student. The first would count as a request for information (Do you need help?), where the second is more likely a request for action (I don't understand, help me.) The importance of context is what makes recognizing illocutionary force difficult for an automated system.

When dealing with data from a platform such as Twitter, however, this becomes less of an issue, since in one-to-many communication every member of the audience is addressed relatively equally and the general context is the same for each tweet (Twitter). This exploratory study aims to find if indirect SA's play a role in mis-/disinformation tweets and if so, if there are linguistic features identifiable that can capture indirect SA's.

2 METHODS

Using the dataset compiled for the MediaEval 2020 FakeNews Task [3], a qualitative analysis of 130 tweets was performed, coding for direct and indirect SA's in Atlas.ti (version 8.4.5). The coding process was iterative and additional codes were added based on patterns found in the data, such as recurring linguistic features, like coherence markers (e.g., 'no meaningful connectives') and certain SA's (e.g., 'citing'). Some codes were mutually exclusive (such as 'Indirect SA: None' with other Indirect SA's), where other codes were not (e.g., tweets were often coded for multiple direct SA's, see Example 1). Only tweets supporting a conspiracy (e.g., corona is a coverup for 5G deaths or 5G causes corona) were included in the result section, as tweets with a different stance were rarely found (n = 7). Furthermore, only tweets that were *not* part of a thread were analyzed, to ensure that the context did not differ in regard to one-to-many vs. one-to-one interaction. The data was coded iteratively until a saturation point was reached within these inclusion criteria (n = 90). The final code list can be found in Appendix 1. Indirect SA's were only coded for when these were the primary communicative force. For example, the tweet below was coded for 'Indirect SA: concluding', as the implied meaning is most likely the primary meaning. Without the implication the assertions made are simply loose statements.

Anyone else curious about the majority of deaths in china seem to be the same areas they rolled out their stand alone 5G just a couple months ago. Verry few deaths being reported in other areas in comparison. #5G #CoronavirusOutbreak #COVID19 #5gamechanger
(Example 1)

The user asks a question in the first sentence, evidenced by the syntactic structure of the sentence ('direct SA: asking'), even though they did not use punctuation. This is followed by two assertions ('direct SA: asserting'). The intended relation between the three sentences is not made explicit through coherence markers such as meaningful connectives ('no meaningful connectives'). The last sentence does have a lexical cue phrase (phrases that show the relation between sentences or the attitude of the speaker, e.g., 'in my opinion' or, in this tweet, 'in comparison'), that shows the relation between areas the user wishes to point out. As the causal relation is not made explicit, the act of concluding that these assertions are causally related is an indirect SA.

3 RESULTS

Users seem to employ a couple of strategies to avoid outright claiming there is a conspiracy, often communicating through

implication (see Table 2). First, they often omit connectives, leaving the relation between sentences implicit (e.g., Example 1 and Table 3). Second, they tend to cite others (Table 2), such as citing the headline of an article they then link to, or use other means to put a middleman between themselves and what is said, as can be seen in Example 2:

I've been reading a few posts from ppl I know about how the #CoronavirusOutbreak is because of 5G trials and Wuhan was the place that first rolled this out and therefore is seen to be used as a biological warfare weapon... 🤔 #coronavirus

(Example 2)

In this tweet, connectives and lexical cue phrases (because of, therefore) are used to make author intent clear, but the user puts the words in the mouths of 'ppl I know'.

Table 1: Frequency of most used direct SA's

Code (not mutually exclusive)	Number of tweets (n = 90)	Percentage of the dataset
Direct SA:		
- asking	28	32.2%
- asserting	69	76.7%
- citing	24	26.6%
- describing	29	32.2%

Table 2: Distribution of indirect SA's (suggesting, concluding, inviting, describing)

Code	Number of tweets (n=90)	Percentage of the dataset
Indirect SA	59	65.6%
No indirect SA	31	34.4%

Table 3: Distribution of linguistic features

Code (not mutually exclusive)	Number of tweets (n = 90)	Percentage of the dataset
No meaningful connectives	52	57.7%
Lexical cue phrases:		
- Opinion	26	28.9%
- Relation	19	21.1%
Emoji use	12	13.3%

Third, users employ ambiguous phrasing. This can be seen in Example 1, where the question is not clearly stated, since a question mark is omitted and instead an assertion follows immediately. It is thus phrased initially as a question, but the question itself does not seem of much importance. In Example 2, this strategy can also be seen in '🤔'. It is left to the reader to infer what the target of the emoji is; is it the entire prior statement (how awkward that people I know say these things) or only the part describing a possible relation between 5G and corona (there might be a relation between corona and 5G)? Depending on the interpretation, the meaning of the tweet changes completely, even with regard to the user's stance.

4 DISCUSSION

This study found that there are linguistic cues identifiable that capture indirect SA's, such as omission of certain connectives and lexical cue phrases. A possible explanation for these findings is that users might (subconsciously) try to circumvent the *forewarning effect*. This effect has been studied extensively in psychology and suggests that forewarning is a factor that causes resistance to persuasion [4]. Using connectives and lexical cue phrases helps reader comprehension in informative texts, but in persuasive texts can build up the reader's resistance, since they recognize more easily when they are being persuaded [4]. Apart from the omission of connectives and lexical cue phrases, the ambiguity of certain tweets also points to this explanation.

Alternatively, the omission of certain words might also be a result of the affordances [5] of Twitter, where the limited characters per tweet might incentivize users to leave out words they deem unnecessary. However, this does not seem to be the case, as one would expect that emojis would be used quite often, since they leave room for ambiguity while simultaneously only taking up one character space. As seen in Table 1, this is not the case. Additionally, emojis are often used decoratively as well, such as arrows or bullet points, without making use of their potential for ambiguity. Furthermore, citing linked article titles is an interesting practice with these affordances in mind, as it leaves little room for the user's own view. The limited space Twitter affords gives weight to the chosen citation, which in turn creates the implication that the citation is important/true/relevant.

Lastly, leaving things ambiguous and citing others, could point to a user orientation to distance themselves from conspiracy-thinking – a way for users to keep plausible deniability for their support of what is said. It should be noted though, that grammatical or punctuation ambiguity might also be a result of users not being native English-speakers or simply inattentiveness or oversight by the user.

Similar to [6]'s findings on Indonesian hoax data, I found that SA Theory can be useful to analyze mis-/disinformation online data. Where they focused on direct SA's, finding that assertive, directive and expressive SA's were most common. In this study, indirect SA were also considered, showing that communication is often indirect in mis-/disinformation tweets. Future research could compare these findings to non-conspiracy-tweets, to shed some light on which explanation provided here is more plausible and to show if the found linguistic features might aid in distinguishing information from mis-/disinformation. It would also be useful to see if machine-learning might have already been able to distinguish between the two regardless of picking up on the indirect SA's.

ACKNOWLEDGMENTS

I thank Martha Larson for her input and critical eye. I also thank John Keates and Zhengyu Zhao for help with pre-processing.

REFERENCES

- [1] J. R. Searle, 1965. What is a Speech Act? In M. Black (Ed.), *Philosophy in America*, pp. 221-239. Allen and Unwin.
- [2] M. L. Geis, 1995. *Speech acts and conversational interaction*. Cambridge University Press.
- [3] K. Pogorelov, D. T. Schroeder, L. Burchard, J. Moe, S. Brenner, P. Filkukova and J. Langguth 2020. FakeNews: Corona Virus and 5G Conspiracy Task at MediaEval 2020. In *Working Notes Proceedings of the MediaEval 2020 Workshop*.
- [4] J. Kamalski, L. Lentz, T. Sanders and R. A. Zwaan, 2008. The Forewarning Effect of Coherence Markers in Persuasive Discourse: Evidence from Persuasion and Processing. *Discourse Processes*, 45(6), 545-579.
- [5] I. Hutchby, 2001. Technologies, Texts and Affordances. *Sociology*, 35(2), 441-456.
- [6] A. A. Iswara and K. A. Bisena, 2020. Manipulation and Persuasion through Language Features in Fake News. *RETORIKA: Jurnal Ilmu Bahasa*, 6, 26-32.