

Review of: "Automatic Content Analysis Systems: Detecting Disinformation in Social Networks"

Muhammad Asfand-E-Yar¹

¹ Bahria University

Potential competing interests: No potential competing interests to declare.

- The paper titled "Automatic Content Analysis Systems: Detecting Disinformation in Social Networks" explores the application of computational linguistics and machine learning methods to detect disinformation on social media platforms. It outlines various techniques, including text preprocessing, feature extraction, and classification algorithms, to identify fake news. The study highlights the effectiveness of these methods and the potential of advanced models like GPT-4 in improving detection accuracy.
- As the title discusses disinformation in social networks, unlike the discussion in the article, is diverting to fake news and spam configuration.

Disinformation means misleading or biased information, manipulated facts, i.e., propaganda.

While fake news means purposefully crafted sensational, emotionally charged, misleading, or totally fabricated information that mimics the form of mainstream news.

- The authors are mixing up the two different types of misleading information, which are manipulated in different ways, i.e., disinformation and fake news. Scientifically, both have different meanings, and the concepts of the words are different.
- In the keywords, again a mix of words is used; for example, propaganda, disinformation, and fake. All these three have different concepts in the research area.
- In the section "Methods and Tools," the authors use a statement that "The problem of detecting fake sources of information can be similar to the tasks of detecting spam," which is totally controversial in detecting spam information and fake news. The methods and techniques of NLP are totally different in detecting spam and/or fake news.
- The spam classification of data and its detection is totally different. Therefore, it cannot be used to detect fake sources of information, as claimed by the authors, i.e., "This approach can be adapted to detect fake sources of information, ..."
- The statement is not clear what the authors want to discuss. "Once the features are defined, they can be used for classification using various machine-learning techniques that involve training with a teacher." Who is the teacher?
- The statement is non-scientific and not clear what the authors want to discuss. Is this statement required to be written in the scientific journal? "Thanks to Char N-grams, you can effectively analyze text regardless of its language and vocabulary, making it a universal tool for processing text information."
- The authors have used a number of text analysis techniques, which also includes classification techniques, for example, TF-IDF, Bag of words, NB classifier, SVM, KNN, POS, probabilistic context free grammar, N-gram, word2Vec, LSTM, Bert, and ELMo. I think he is missing RNN, Random Forest, logistic regression, Binary Tree, and

many more.

Why are all of them discussed? What is the reason to discuss them? Are they all used in training? If yes, then how?

- It is not clear what is meant by "expert fact-checking" and "crowdsourced fact-checking". Where are the technical proofs, i.e., how the experiments are carried out according to these two categories?
- While the paper mentions the dataset used for the experiments, it lacks a detailed description of its sources, characteristics, and any potential biases.
- The results are based on a specific dataset, and there is limited discussion on the generalizability of the findings to other datasets or real-world scenarios.
- The paper could benefit from a more detailed comparative analysis with other state-of-the-art methods to highlight the advantages and limitations of the proposed approach.
- The paper does not address the ethical implications of using machine learning for disinformation detection, such as potential biases in the models and the impact on freedom of expression.

Overall, the paper lacks research, i.e., the experimental part needs discussion according to the framework architecture of the working model. The results are only the display of the dataset in different formats; still, the graphs are generated via a Python library that discusses the dataset presentations.

With some improvements, particularly in the areas of dataset description, comparative analysis, and ethical scientific considerations, it would be a strong candidate for publication in the journal.

Ethical considerations means to address the ethical implications of the research, including potential biases in the models and the impact on user privacy and freedom of expression.