

Is traditional natural language processing ‘dead’?

Saranzaya Magsarjav ^a 

^a *School of Computer and Mathematical Sciences, University of Adelaide, Adelaide, 5005, Australia*

Email: saranzaya.magsarjav@adelaide.edu.au

Abstract: Over the last few years, there has been a huge surge in Large Language Models in Natural Language Processing (NLP) tasks. They have been shown to outperform many traditional models in tasks. This has pushed traditional natural language processing out of the spotlight to a point where they have been claimed to be irrelevant and useless, in some extreme claims, ‘dead’. While these claims are not unfounded, they lack some deeper understanding of when and where each model is useful. Just like any other models and techniques, large language models have their flaws and drawbacks.

One of the reasons why LLMs work so well and have taken off is due to the extensive amount of text data available now, with the relatively free access of the internet. While this is one of the reasons why it works so well, it is also one of the drawbacks. As it is trained on online content, there is a huge amount of biased data that is present. It has been shown that LLMs amplify and relay this in the outputs, perpetuating these biases further.

While the performances of LLMs are good, there are huge problems in output accuracy due to hallucinations. While traditional natural language processing will output information based on the data, and it is the user's place to interpret this. While LLMs will output information as if it is factual and will also give it context, whether this context exists or not.

All of these problems are encompassed by the main problem that LLMs are a black box (Hadi et al. 2023). There is a lack of information on the architecture and the training data. As you cannot ‘look under the hood’ of these models, it is hard to explain why they create certain outputs. This causes problems in reproducibility and reliability. There is no guarantee that another researcher can recreate these outputs, and sometimes even the researcher themselves cannot recreate the outputs.

While traditional Natural Language Processing may perform worse compared to LLMs, they have a distinct advantage in that they are vastly more explainable and do not need as many resources. When it comes to high-risk decision-making, LLMs may give more ‘accurate’ information, but when it cannot be backed up with specific decisions the model has made, it makes it less trustworthy. In fields such as forensics and health sciences, there is a high risk of failure; there needs to be a greater amount of accountability in the model's choices. This is where traditional natural language processing shines. Traditional Natural Language Processing allows us to point to specific features in the dataset or model that help explain the decisions.

This talk will demonstrate some of these inconsistencies present in LLMs and showcase when traditional NLP is appropriate and is not ‘dead’ in text-based research.

REFERENCES

Hadi MU, Qureshi R, Shah A, Irfan M, Zafar A, Shaikh MB, Akhtar N, Wu J and Mirjalili S (2023) Large language models: a comprehensive survey of its applications, challenges, limitations, and future prospects. *Authorea preprints*, 1(3), pp.1-26.

Keywords: *Natural Language Processing, Large Language Models*