

REVIEWER'S REPORT

Manuscript No.: JNHST-028

Title: Cognitive Divergence as a Predictor of LLM Output Comprehensibility: Extending Relative-to-Human Benchmarks Beyond Translation Tasks

Recommendation:

Accept as it is
Accept after minor revision.....
Accept after major revision
Do not accept (*Reasons below*)

Rating	Excel.	Good	Fair	Poor
Originality	...			
Techn. Quality		...		
Clarity	...			
Significance	...			

Reviewer's ID: JNJST-06

Detailed Reviewer's Report

The paper titled “**Cognitive Divergence as a Predictor of LLM Output Comprehensibility: Extending Relative-to-Human Benchmarks Beyond Translation Tasks**” presents a significant and timely contribution to the field of natural language processing by addressing a critical gap in the evaluation of large language model (LLM) outputs. The central argument of the study is that existing evaluation metrics such as BLEU and ROUGE are insufficient for assessing true comprehensibility, as they rely heavily on surface-level textual similarity rather than deeper cognitive alignment with human understanding. The authors introduce the concept of Cognitive Divergence, defined as the multidimensional difference between machine-generated and human-generated text across syntactic, semantic, and discourse levels. This shift from absolute to relative, human-centered benchmarking forms the conceptual backbone of the paper and reflects a progressive move toward more meaningful AI evaluation frameworks.

The literature review is comprehensive and well-structured, critically examining both traditional and neural evaluation metrics. It effectively highlights the limitations of existing approaches, particularly their inability to capture cognitive and discourse-level aspects of language comprehension. The integration of cognitive and psycholinguistic theories, such as Cognitive Load Theory and Dependency Length Minimization, strengthens the theoretical grounding of the study. By linking these theories to NLP evaluation, the authors successfully establish an interdisciplinary framework that bridges artificial intelligence with human cognitive processing, which is one of the key strengths of the paper. Methodologically, the study adopts a mixed-methods approach, combining quantitative statistical analysis with qualitative discourse evaluation. The use of multiple datasets including translation, summarization, and question-answering tasks enhances the robustness and generalizability of the findings. The operationalization of Cognitive Divergence into syntactic, semantic, and discourse components is particularly noteworthy, as it allows for a nuanced and multidimensional analysis. Furthermore, the inclusion of human evaluation scores alongside automated metrics adds credibility and validity to the results, ensuring that the findings are grounded in real human judgments of comprehensibility.

The results of the study strongly support the proposed hypothesis, demonstrating that Cognitive Divergence is a more reliable predictor of comprehensibility than traditional and even neural evaluation metrics. The negative correlation between divergence and human comprehension scores is clearly established through both correlation and regression analyses. The qualitative findings further enrich the discussion by identifying specific issues in LLM outputs, such as syntactic misalignment, semantic ambiguity, and discourse-level inconsistencies. These insights not only validate the Cognitive Alignment Hypothesis but also provide practical implications for improving LLM design and evaluation. In terms of contributions, the paper makes a valuable addition to the growing discourse on human-centered AI by proposing a holistic and cognitively informed evaluation framework. Its implications extend beyond

REVIEWER'S REPORT

NLP research to areas such as education, human-computer interaction, and AI system design. However, the study could be further strengthened by providing more detail on the computational implementation of the Cognitive Divergence Score and addressing potential scalability challenges in real-world applications. Additionally, while the interdisciplinary approach is a strength, some sections could benefit from clearer explanations to make the concepts more accessible to readers from non-technical backgrounds.

Overall, the paper is well-argued, methodologically sound, and theoretically rich. It successfully challenges existing paradigms in LLM evaluation and offers a promising direction for future research. By emphasizing cognitive alignment as a key determinant of comprehensibility, the study lays the groundwork for more human-centric and effective AI communication systems.