

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

Abstract

This research investigates the role of Cognitive Divergence as a critical factor in determining the comprehensibility of Large Language Models' (LLMs) outputs. Despite recent breakthroughs in artificial intelligence that have greatly enhanced the readability and grammar of machine-generated text, there is uncertainty about the comprehensibility of such language for human readers. Established evaluation measures such as BLEU, ROUGE, and other textual similarity measures primarily focus on surface-level textual similarity and may not reflect deeper cognitive compatibility between human and machine-generated language. To overcome this challenge, this study proposes a benchmarking framework that considers LLM outputs in relation to human-generated text under similar circumstances. Cognitive Divergence is defined as a multi-dimensional difference, including syntax, semantics, and discourse structure. The research uses a multi-pronged approach, integrating quantitative and qualitative discourse analysis in various natural language processing tasks such as translation, summarization and question answering. The results show that Cognitive Divergence is a better predictor of comprehensibility than conventional evaluation measures. Texts with lower divergence scores are perceived as more comprehensible than others, even when factors such as grammatical accuracy and vocabulary complexity are accounted for. Moreover, the findings show this approach is not only suitable for the task of translation but can also be applied to a range of text generation scenarios. This research adds to the debate about human-centered AI evaluation by introducing a holistic approach that combines cognitive science with relative benchmarking. It also underscores the need to align machine language with human cognitive processes to improve human-machine interaction. This study has implications for the development of future evaluation metrics, creation of more human-centred language models, and enhancement of AI communication technologies.

Keywords: Cognitive Divergence, Cognitive Load Theory, Dependency Length Minimization, Large Language Models, Psycholinguistic Theories

Introduction

Large Language Models (LLMs) represent a significant breakthrough in natural language processing (NLP), with impressive abilities to generate human-like text for a variety of applications, such as translation, summarization, conversational agents and even writing papers (Brown et al., 2020; Achiam et. al., 2023). These models, which are based on transformer networks, are trained on large data sets to generate coherent and relevant text (Vaswani et al., 2017). However, a key question remains: how do we assess the readability of these outputs for human consumption?

Classic evaluation metrics, like BLEU and ROUGE, have been popular for evaluating machine-produced text, using the overlap between generated text and reference text as a metric (Papineni et al., 2002; Lin, 2004). These metrics offer a formalized way of evaluation, but they lack consideration for deeper aspects of meaning, coherence, and human understanding. Specifically, they are based on the assumption that more similarity to a reference text led to better quality, which is not the case in many open-ended language generation tasks (Callison-Burch et al., 2006). In response, neural evaluation metrics have been developed, such as BERT Score and COMET, which use contextual embeddings to compute semantic similarity (Zhang et al., 2019; Rei et al., 2020).

These methods have shown better alignment with human evaluation; however, they are still largely black boxes and do not explicitly include models of human cognitive processing. This prevents them from fully accounting for why some texts are easier to comprehend than others. In recent years, research has started to focus on considering cognitive and psycholinguistic aspects in language assessment. Theories like Cognitive Load Theory propose that comprehension is affected by the cognitive load associated with processing information (Sweller, 1988). Linguistic principles such as Dependency Length Minimization suggest humans also favour syntactic structures that ease comprehension (Futrell et al., 2015). But empirical evidence indicates that absolute measures of complexity are not enough to explain comprehensibility (Gibson, 1998; Touvron et al., 2023).

This paper addresses these problems by proposing Cognitive Divergence as the extent to which machine-generated text diverges from human cognitive and linguistic processing patterns. Instead of using absolute measures, Cognitive Divergence uses a benchmarking approach relative to human performance, by comparing the output of LLMs with that of humans in identical tasks. This approach reframes the problem from linguistic to cognitive similarity.

The key hypothesis of this paper is that text comprehensibility is not merely a function of linguistic correctness and complexity but rather alignment with human expectations and cognition. Thus, assessing LLM generated output needs to consider these cognitive aspects. Moreover, while past work has mostly focused on such concepts in the context of machine translation, this paper applies these measures to different NLP tasks such as summarization and question answering. This will enable the establishment of Cognitive Divergence as a universal measure of comprehensibility.

The rest of this paper is organized as follows: the literature review identifies current evaluation approaches and cognitive theories; the theoretical framework outlines the concept of Cognitive Divergence; the methodology section gives an overview of the research design; and the results and discussion section presents the findings and their discussion.

Literature Review

Traditional Evaluation Metrics

Automatic evaluation metrics such as BLEU, ROUGE and METEOR have been used to assess the quality of machine-generated text. BLEU (Bilingual Evaluation Understudy) counts the overlap of n-grams between a candidate translation and one or more reference translations (Papineni et al., 2002). ROUGE, on the other hand, measures recall-based overlap, often used in summarization (Lin, 2004), and METEOR accounts for synonymy and stemming to enhance the evaluation process (Banerjee & Lavie, 2005). However, these metrics have been shown to miss semantic and discourse level aspects of quality.

They rely on the assumption that similarity in words is synonymous with quality, which can result in misleading results, especially in tasks with multiple valid solutions (Callison-Burch et al., 2006). Moreover, these measures are not sensitive to different phrasings that humans might find acceptable. Text comprehensibility has also been measured using readability formulas like the Flesch Reading Ease score. But these metrics are largely based on sentence length and word complexity, and overlook more complex cognitive and contextual considerations (Pitler & Nenkova, 2008). Thus, they offer a limited perspective on text comprehensibility.

Neural Evaluation Metrics

Neural-based metrics have been proposed to address shortcomings of traditional metrics. The BERT Score assesses text quality by measuring embeddings from pre-trained language models (Zhang et al., 2019). COMET employs neural networks trained with human judgement data to assess translation quality (Rei et al., 2020). Both metrics show better agreement with human judgements, especially in machine translation. But they have limitations. First, they are black-box models, with little explanation of how scores are derived (Celikyilmaz et al., 2020). Second, they do not explicitly leverage cognitive principles, and it is hard to interpret their relationship with human understanding. Third, neural metrics still require reference-based comparisons, which can be absent or inappropriate in some tasks, such as generating dialogues and writing stories. This makes it apparent that we need task-independent and cognitively-based frameworks for evaluation.

Cognitive and Psycholinguistic Theories

Cognitive theories can tell us a lot about language processing. Cognitive Load Theory proposes that ease of comprehension is affected by the processing demands of linguistic information (Sweller, 1988). Passages with high cognitive load are harder to comprehend. Dependency Length Minimization (DLM) proposes that languages minimize the length of dependencies to ease comprehension (Futrell et al., 2015). Likewise, Gibson (1998) claims that comprehension is influenced by linguistic complexity via memory limitations. But such models tend to only consider absolute complexity and ignore the variability in language use. Research shows that texts with comparable levels of complexity can vary in comprehensibility, implying that other factors such as familiarity and expectations are important.

Discourse and Coherence

In addition to sentence-level aspects, discourse features are essential for comprehension. Halliday and Hasan (2014) stress the cohesive function of text, connecting various elements, while Rhetorical Structure Theory (Mann, 1987) points to the importance of relations among discourse units. Kintsch (1998) argue comprehension involves building a mental model of the text, which requires coherence and logical structure. Discourse errors can have a major impact on comprehension, even if sentences are grammatically accurate. This research implies that when assessing LLM responses, discourse features - which are often ignored by existing metrics - must be considered.

Cognitive Divergence

A New Approach A recent study has proposed Cognitive Divergence as a metric to evaluate the differences between human and LLM outputs (Liu, 2025; Lu et al., 2023). This moves away from absolute measures to relative ones, with a focus on human cognitive patterns. Cognitive Divergence measures variations across syntactic, semantic and discourse levels. Early research suggests a correlation between higher divergence and lower readability, suggesting its potential as a predictor. Furthermore, relative-to-human comparisons are in line with recent shifts in AI evaluation towards human-centered methods (Liang et al., 2022).

Theoretical Framework

This research is based on a multidisciplinary theoretical framework combining cognitive psychology, linguistics and artificial intelligence evaluation. The notion of Cognitive Divergence is situated at the juncture of these fields, building on existing theories of human comprehension and applying them to the assessment of AI-generated text.

Cognitive Load Theory

Cognitive Load Theory (CLT) suggests that humans have a limited capacity for working memory and that understanding is based on the effective processing of information (Sweller, 1988). Foreigner texts that overload intrinsic and/or extraneous cognitive load stop comprehension, even if they are grammatically correct. For LLM-generated texts, even grammatically correct sentences may not be easy to understand if they are cognitively demanding. Cognitive load increases when language diverges from human preferences in cognitive divergence idea. If LLM output deviates from human norms, readers have to work harder to make sense of it, leading to lower comprehensibility.

Dependency Length Minimization (DLM)

DLM theory suggests languages strive to keep syntactically dependent words closer together to ease comprehension (Futrell et al., 2015). Minimizing dependency length decreases working memory load and enhances reading. But LLM generated text may not always observe these principles. Cognitive Divergence quantifies such deviations in terms of the differences in the

147 syntax of human and machine texts. This finding is consistent with research showing processing
148 efficiency is crucial for comprehension (Gibson, 1998).

149 **Processing and Coherence**

150 Readers don't just process and understand sentences but also build a meaningful representation of
151 the text as a whole (Kintsch, 1998). Discourse theories highlight the roles of coherence, cohesion
152 and logical organization (Halliday & Hasan, 2014; Mann, 1987). Cognitive Divergence builds on
153 these notions and assesses the extent to which LLM-generated texts maintain discourse-level
154 congruence with human texts. Inconsistencies such as topic shifts and poor logical connections
155 are seen as signs of divergence.

156 **Relative-to-Human Benchmarking**

157 Most evaluation metrics use absolute measures, but recent research has argued for the need for
158 relative-to-human evaluation (Freitag et al., 2021). Relative-to-human assessments includes
159 comparison of machine output to text produced by humans in the same task. This study
160 incorporates this method as part of Cognitive Divergence. Using human output as the standard,
161 the approach takes into account the variation in acceptable language usage and offers a more
162 meaningful quality assessment.

163 **Cognitive Alignment Hypothesis**

164 Drawing on the above theories, this study puts forward the Cognitive Alignment Hypothesis: The
165 extent to which machine-generated text aligns with human cognitive-linguistic processes affects
166 its comprehensibility. This hypothesis underpins the empirical studies in this work.

167 **Methodology**

168 **Research Design**

169 The research adopts a quantitative- qualitative mixed-methods approach. The aim is to
170 investigate the association between Cognitive Divergence and comprehensibility in various NLP
171 tasks.

172 **Data Sources**

173 We used three types of data:

174 1. Machine Translation Data Created from previous research's standard benchmark (Hassan et
175 al., 2018).

176 2. Summarization Data Based on news summarization data (Narayan et al., 2018).

177 3. Question Answering Data Adapted from SQuAD (Rajpurkar et al., 2016). For each of the
178 datasets, we gathered both human and LLM responses for comparison.

179 **Operationalizing Cognitive Divergence**

Cognitive Divergence was operationalized in three ways:

- Syntactic Divergence: Variations in syntax and dependency
- Semantic Divergence: Differences in semantic representation
- Discourse Divergence: Coherence and discourse structure

These were calculated using both automated and manual methods.

Evaluation Metrics

We used the following evaluation metrics:

- Cognitive Divergence Score (CDS) (new)
- COMET Score (Rei et al., 2020)
- BERTScore (Zhang et al., 2019)
- Human Evaluation Scores (Celikyilma et al., 2020) Humans scored outputs for comprehensibility on a 5-point scale

Data Analysis Techniques

The following analytical methods were applied:

- Correlation Analysis: To assess the correlation between CDS and human scores
- Regression Analysis: To assess predictive power
- Discussion Analysis: To detect qualitative themes Data were statistically analyzed to determine p-values and effect sizes.

Reliability and Validity

Reliability was assured by using multiple human annotators, for which inter-rater agreement was measured. Validity was ensured by comparing CDS to other metrics and human ratings. Step 6:

Results and Discussion

Quantitative Findings

Our findings show a strong negative relationship between Cognitive Divergence and comprehensibility. Increasing divergence is associated with decreasing human scores. The traditional metrics (BLEU and ROUGE) show low correlation, as previously reported (Papineni et al., 2002; Lin, 2004). Neural metrics like COMET and BERTScore have higher moderate correlation but are still inferior to Cognitive Divergence (Rei et al., 2020; Zhang et al., 2019).

Regression Analysis

Linear regression shows that Cognitive Divergence is a significant predictor of comprehension, better than all baseline metrics. CDS increases the explanatory power.

Qualitative Analysis

Qualitative analysis reveals a number of insights:

- Syntactic Misalignment: LLM responses have unnatural sentence structure
 - Semantic Ambiguity: Shifts in meaning create ambiguity
 - Discursive Disconnect: Poor connectives confuse
- These results are consistent with the Cognitive Alignment Hypothesis.

Generalization Across Tasks

Our approach works across:

- Translation tasks (Hassan et al., 2018)
 - Summarization tasks (Maynez et al., 2020)
 - Question answering tasks (Rajpurkar et al., 2016)
- This shows its consistency and versatility.

Theoretical Implications

The findings indicate that comprehensibility is a cognitive rather than a linguistic concept. This calls for a rethink of current levels of evaluation, and the need for human-centred metrics.

Implications

Implications for NLP

The current research reveals the lack of effectiveness in conventional evaluation metrics and emphasizes the need for cognitive evaluation (Liang et al., 2022). Researchers should work towards incorporating cognitive considerations into evaluation.

Design Recommendations for AI

AI text generators should be human-friendly. They should be trained to be not only fluent but also cognitively compatible with human readers (Bowman et al., 2022).

Implications for Education

Cognitive Divergence can be applied to assess AI outputs in educational settings, to ensure that they are accurate as well as comprehensible.

Implications for Human-AI Interaction

Better comprehension increases the human trust and acceptance of AI systems, which is crucial in critical areas like healthcare and education.

Conclusion

This research confirms Cognitive Divergence as a valid measure of the comprehensibility of LLM-generated text. It offers a more relevant and human-centred approach by moving away from absolute and towards relative-to-human text metrics. The research shows that adherence to human cognitive and language processing patterns is critical in achieving effective communication. Existing metrics, although valuable, overlook this aspect, suggesting more sophisticated methods of evaluation. Moreover, this research generalizes the notion of Cognitive Divergence beyond translation, showing its use in other NLP applications. Such generalization is a key advancement. It would be interesting for future work to develop automated approaches to compute Cognitive Divergence and to use it in real-time evaluation applications. Finally, incorporating this approach into training could enhance LLM performance. Ultimately, this study adds to the emerging research on human-centred AI by suggesting a cohesive, cognitively-driven evaluation framework. Its focus on human understanding lays the foundation for more efficient and reliable language tools.

References

- Banerjee, S., & Lavie, A. (2005, June). METEOR: An automatic metric for MT evaluation with improved correlation with human judgments. In *Proceedings of the acl workshop on intrinsic and extrinsic evaluation measures for machine translation and/or summarization* (pp. 65-72).
- Brown, T., Mann, B., Ryder, N., Subbiah, M., Kaplan, J. D., Dhariwal, P., & Amodei, D. (2020). Language models are few-shot learners. *Advances in neural information processing systems*, 33, 1877-1901.
- Bowman, S. R., Hyun, J., Perez, E., Chen, E., Pettit, C., Heiner, S., & Kaplan, J. (2022). Measuring progress on scalable oversight for large language models. *arXiv preprint arXiv:2211.03540*.
<https://doi.org/10.48550/arXiv.2211.03540>
- Callison-Burch, C., Osborne, M., & Koehn, P. (2006, April). Re-evaluating the role of BLEU in machine translation research. In *11th conference of the european chapter of the association for computational linguistics* (pp. 249-256).
- Freitag, M., Foster, G., Grangier, D., Ratnakar, V., Tan, Q., & Macherey, W. (2021). Experts, errors, and context: A large-scale study of human evaluation for machine translation. *Transactions of the Association for Computational Linguistics*, 9, 1460-1474.
https://doi.org/10.1162/tacl_a_00437
- Futrell, R., Mahowald, K., & Gibson, E. (2015). Large-scale evidence of dependency length minimization in 37 languages. *Proceedings of the National Academy of Sciences*, 112(33), 10336-10341. <https://doi.org/10.1073/pnas.1502134112>

274 Celikyilmaz, A., Clark, E., & Gao, J. (2020). Evaluation of text generation: A survey. *arXiv*
 275 *preprint arXiv:2006.14799*. <https://doi.org/10.48550/arXiv.2006.14799>

276 Gibson, E. (1998). Linguistic complexity: Locality of syntactic dependencies. *Cognition*, 68(1),
 277 1-76. [https://doi.org/10.1016/S0010-0277\(98\)00034-1](https://doi.org/10.1016/S0010-0277(98)00034-1)

278 Halliday, M. A. K., & Hasan, R. (2014). *Cohesion in english*. Routledge.

279 Hassan, H., Aue, A., Chen, C., Chowdhary, V., Clark, J., Federmann, C., & Zhou, M. (2018).
 280 Achieving human parity on automatic chinese to english news translation. *arXiv preprint*
 281 *arXiv:1803.05567*. <https://doi.org/10.48550/arXiv.1803.05567>

282 Kintsch, W. (1998). *Comprehension: A paradigm for cognition*. Cambridge university press.

283 Lu, Q., Ding, L., Xie, L., Zhang, K., Wong, D. F., & Tao, D. (2023, July). Toward human-like
 284 evaluation for natural language generation with error analysis. In *Proceedings of the 61st Annual*
 285 *Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)* (pp. 5892-
 286 5907). <https://doi.org/10.18653/v1/2023.acl-long.324>

287 Liang, P., Bommasani, R., Lee, T., Tsipras, D., Soylu, D., Yasunaga, M., & Koreeda, Y. (2022).
 288 Holistic evaluation of language models. *arXiv preprint arXiv:2211.09110*.
 289 <https://doi.org/10.48550/arXiv.2211.09110>

290 Lin, C. Y. (2004, July). Rouge: A package for automatic evaluation of summaries. In *Text*
 291 *summarization branches out* (pp. 74-81).

292 Mann, W. (1987). Rhetorical structure theory: Toward a functional theory of text
 293 organization. *Text*, 8(3), 167-182.

294 Maynez, J., Narayan, S., Bohnet, B., & McDonald, R. (2020, July). On faithfulness and
 295 factuality in abstractive summarization. In *Proceedings of the 58th annual meeting of the*
 296 *association for computational linguistics* (pp. 1906-1919). [https://doi.org/10.18653/v1/2020.acl-](https://doi.org/10.18653/v1/2020.acl-main.173)
 297 [main.173](https://doi.org/10.18653/v1/2020.acl-main.173)

298 Narayan, S., Cohen, S. B., & Lapata, M. (2018). Don't give me the details, just the
 299 summary. *Topic-aware convolutional neural networks for extreme summarization*. *ArXiv*
 300 *abs/1808.08745*.

301 Achiam, J., Adler, S., Agarwal, S., Ahmad, L., Akkaya, I., Aleman, F. L., & McGrew, B. (2023).
 302 Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*.
 303 <https://doi.org/10.48550/arXiv.2303.08774>

304 Papineni, K., Roukos, S., Ward, T., & Zhu, W. J. (2002, July). Bleu: a method for automatic
 305 evaluation of machine translation. In *Proceedings of the 40th annual meeting of the Association*
 306 *for Computational Linguistics* (pp. 311-318).

- 307 Pitler, E., &Nenkova, A. (2008, October). Revisiting readability: A unified framework for
308 predicting text quality. In *Proceedings of the 2008 conference on empirical methods in natural*
309 *language processing* (pp. 186-195).
- 310 Rajpurkar, P., Zhang, J., Lopyrev, K., & Liang, P. (2016, November). Squad: 100,000+ questions
311 for machine comprehension of text. In *Proceedings of the 2016 conference on empirical methods*
312 *in natural language processing* (pp. 2383-2392).
- 313 Rei, R., Stewart, C., Farinha, A. C., & Lavie, A. (2020, November). COMET: A neural
314 framework for MT evaluation. In *Proceedings of the 2020 conference on empirical methods in*
315 *natural language processing (emnlp)* (pp. 2685-2702).
- 316 Liu, J. (2025). Large Language Models as Annotators for Multi-source, Multi-class Mental-
317 Health Text: Agreement, Uncertainty, and Learnability. *Multi-class Mental-Health Text:*
318 *Agreement, Uncertainty, and Learnability* (December 24, 2025).
319 <http://dx.doi.org/10.2139/ssrn.5964738>
- 320 Sweller, J. (1988). Cognitive load during problem solving: Effects on learning. *Cognitive*
321 *science*, 12(2), 257-285. [https://doi.org/10.1016/0364-0213\(88\)90023-7](https://doi.org/10.1016/0364-0213(88)90023-7)
- 322 Touvron, H., Lavril, T., Izacard, G., Martinet, X., Lachaux, M. A., Lacroix, T., & Lample, G.
323 (2023). Llama: Open and efficient foundation language models. *arXiv preprint*
324 *arXiv:2302.13971*.
- 325 Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., &Polosukhin, I.
326 (2017). Attention is all you need. *Advances in neural information processing systems*, 30.
- 327 Zhang, T., Kishore, V., Wu, F., Weinberger, K. Q., & Artzi, Y. (2019). Bertscore: Evaluating text
328 generation with bert. *arXiv preprint*
329 *arXiv:1904.09675*.<https://doi.org/10.48550/arXiv.1904.09675>

330