

Performance Evaluation of Large Language Models: A Comprehensive Review

Divyakant Meva*

Associate Professor, Faculty of Computer Applications,
Marwadi University,
Rajkot, Gujarat 360003, INDIA

divyakant.meva@marwadieducation.edu.in

Hirenkumar Kukadiya

Assistant Professor, Gandhinagar Institute of Computer Science and Applications,
Gandhinagar University,
Gandhinagar, Gujarat 382721, INDIA

hiren.kukadiya@gandhinagaruni.ac.in



Publication History

ManuscriptReference:IRJCS/RS/Vol.12/Issue03/MRCSI0084| Research Article | Open Access | Double-Blind Peer Reviewed
Article ID: IRJCS/RS/Vol.12/Issue03/MRCSI0084

Received: 06, March 2025, Revised: 12, March 2025 Accepted: 21 March 2025 Published Online: 29 March 2025

<http://www.irjcs.com/volumes/Vol12/iss-03/03.MRCSI0084.pdf>

Article Citation: Divyakant,Hirenkumar(2025).Performance Evaluation of Large Language Models:A Comprehensive Review. IRJCS: International Research Journal of Computer Science, Volume 12, Issue 03 of 2025 pages 109-114

doi:> <https://doi.org/10.26562/irjcs.2025.v1203.03>

BibTeX Bhavishya@2025Facial



Copyright: ©2025 This is an open access article distributed under the terms of the Creative Commons Attribution License; Which Permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.

Abstract:The performance evaluation techniques for large language models (LLMs) are thoroughly reviewed in this study. As these models' capabilities and applications continue to develop quickly, establishing robust evaluation frameworks becomes increasingly important. We analyze the evolution of evaluation metrics and benchmarks, from traditional natural language processing assessments to more recent LLM-specific frameworks. The paper includes detailed performance comparisons across leading models including GPT-4, Claude 3, LLaMA 2, Gemini, and PaLM 2 on established benchmarks. We identify key challenges in current evaluation practices, including reliability concerns, emergent capability assessment, and alignment evaluation. Additionally, we examine the multi-dimensional nature of LLM evaluation, encompassing not only task performance but also safety, efficiency, fairness, and environmental impact considerations. The review concludes with recommendations for developing more holistic evaluation methodologies and identifies interesting avenues for further study in this quickly developing discipline.

Keywords:AI Model Assessment, Benchmarking, Computational Efficiency, Comparative Analysis, Performance Evaluation

I. INTRODUCTION

Large Language Models (LLMs) have emerged as transformative technologies in artificial intelligence, demonstrating remarkable capabilities across diverse tasks from text generation to complex reasoning and problem-solving. These models, trained on vast corpora of text data, have grown exponentially in scale and sophistication over recent years, with models like GPT-4 (OpenAI, 2023), Claude 3 (Anthropic, 2023), LLaMA 2 (Meta AI, 2023), Gemini (Google, 2023), and PaLM 2 (Google, 2023) pushing the boundaries of what's possible in natural language processing. As LLMs increasingly power applications across domains from content creation and summarization to coding assistance and decision support—the need for robust, comprehensive evaluation frameworks has become paramount. Traditional metrics designed for earlier NLP systems often fail to capture the nuanced performance of modern LLMs, particularly their emergent capabilities, instruction-following abilities, and alignment with human values. This review systematically examines the current landscape of LLM evaluation, synthesizing findings from recent literature and industry practices. We provide a detailed analysis of evaluation methodologies, benchmarks, and comparative performance across leading models. Additionally, we identify critical gaps in existing approaches and propose directions for developing more holistic evaluation frameworks. The paper is structured as follows: Section 2 outlines the evolution of evaluation metrics from traditional NLP assessments to LLM-specific approaches. Section 3 presents comprehensive performance comparisons across leading models on established benchmarks. Section 4 discusses key challenges in current evaluation practices. Section 5 examines multi-dimensional evaluation frameworks. Section 6 explores future directions and recommendations. Section 7 offers concluding remarks.

II. EVOLUTION OF LLM EVALUATION METRICS

A. Traditional NLP Metrics

Early evaluation of language models relied on metrics designed for specific NLP tasks:

Perplexity: A statistical measure of how well a probability model predicts a sample, with lower values indicating better prediction capability. While useful for comparing language modeling quality, perplexity provides limited insight into practical capabilities for downstream tasks (Radford et al., 2019).

BLEU, ROUGE, and METEOR: Originally developed for machine translation and summarization tasks, these metrics assess n-gram overlap between model outputs and reference texts. Their correlation with human judgment is often limited, particularly for creative or diverse text generation (Callison-Burch et al., 2006).

Accuracy, Precision, Recall, and F1: Standard classification metrics used for tasks like sentiment analysis, named entity recognition, and question answering. While straightforward, these metrics often fail to capture the quality and usefulness of generated text (Wang et al., 2018).

B. LLM-Specific Evaluation Frameworks

As LLMs evolved, more sophisticated evaluation approaches emerged:

Benchmark Suites: Comprehensive collections of tasks designed to evaluate models across multiple dimensions. Notable examples include:

- **GLUE and SuperGLUE:** General Language Understanding Evaluation benchmarks covering various NLU tasks (Wang et al., 2018; Wang et al., 2019).
- **MMLU (Massive Multitask Language Understanding):** Tests knowledge across 57 subjects from elementary to professional levels (Hendrycks et al., 2021).
- **BIG-Bench:** Collaborative benchmark with 204 tasks spanning reasoning, knowledge, and social bias (Srivastava et al., 2022).
- **HumanEval and MBPP:** Benchmarks for code generation capabilities (Chen et al., 2021).

Human Evaluation Frameworks: Structured approaches for collecting human judgments on model outputs:

- **Likert Scale Ratings:** Direct assessment of quality, relevance, or helpfulness (Clark et al., 2021).
- **Pairwise Comparisons:** Human evaluators choose between outputs from different models (Ouyang et al., 2022).
- **Turing Tests:** Evaluating whether model outputs are distinguishable from human-written content (Brown et al., 2020).

Automated Evaluation Systems:

- **LM-Eval Harness:** An open-source framework for standardized evaluation across multiple benchmarks (Gao et al., 2023).
- **HELM (Holistic Evaluation of Language Models):** Standardized evaluation across dimensions including accuracy, calibration, robustness, fairness, and bias (Liang et al., 2022).

III. PERFORMANCE COMPARISON OF LEADING LLMs

A. Benchmark Performance Comparison

Table 1 presents a comprehensive comparison of performance across leading LLMs on established benchmarks. Values represent percentage scores, with higher values indicating better performance.

Table 1: Performance Comparison of Leading LLMs on Key Benchmarks (%)

Model	MMLU	HumanEval	GSM8K	TruthfulQA	HellaSwag	BBH
GPT-4	86.4	67.0	92.0	59.0	95.3	83.1
Claude 3 Opus	86.8	75.2	94.2	64.8	95.9	88.6
Gemini Ultra	90.0	74.4	94.4	59.4	95.9	83.6
LLaMA 2 (70B)	68.9	29.9	56.8	41.2	85.5	51.2
PaLM 2	78.3	43.4	80.7	53.4	93.0	67.3
Falcon (180B)	70.1	31.5	58.4	45.0	87.1	55.7
Mixtral 8x7B	70.6	38.2	74.9	46.2	88.0	57.3

Note: Values compiled from published evaluations and technical reports. Some values are approximated from available literature where exact figures were not reported.

B. Analysis of Performance Trends

Several key trends emerge from the performance comparison:

1. **Scale Advantage:** Larger models consistently outperform smaller ones across most benchmarks, though with diminishing returns beyond certain scale thresholds. The performance gap between frontier models (>100B parameters) and smaller models (7-70B parameters) remains significant, particularly on reasoning-intensive tasks like GSM8K and BBH.
2. **Reasoning Capabilities:** Models exhibit varying strengths in reasoning tasks. Claude 3 Opus and Gemini Ultra demonstrate the strongest performance on mathematical reasoning (GSM8K) and complex reasoning (BBH), while GPT-4 shows comparable but slightly lower performance on these tasks.
3. **Knowledge Breadth:** On MMLU, which tests knowledge across diverse domains, Gemini Ultra achieves the highest score, followed closely by Claude 3 Opus and GPT-4. This suggests comparable knowledge breadth across frontier models.

4. **Coding Proficiency:** Claude 3 Opus leads in code generation capabilities as measured by HumanEval, followed by Gemini Ultra and GPT-4, while other models lag significantly.
5. **Truthfulness:** Claude 3 Opus demonstrates the highest performance on TruthfulQA, suggesting stronger resistance to generating false information compared to other models.

C. Real-World Performance Considerations

Benchmark performance, while informative, does not always translate directly to real-world utility. Several studies have examined performance in applied settings:

- **Professional Domain Applications:** Nori et al. (2023) evaluated LLMs on medical, legal, and financial tasks, finding that frontier models (GPT-4, Claude 3) achieved 80-90% accuracy on professional licensing exam questions but struggled with complex domain-specific reasoning.
- **Tool Use and Planning:** When evaluated on tasks requiring external tool use, planning, and multi-step reasoning, model performance shows higher variance than on static benchmarks (Schick et al., 2023).
- **Instruction Following:** Models exhibit varying capabilities in following complex, multi-part instructions, with instruction-tuned variants showing significant advantages over base models (Wei et al., 2022).

IV. CHALLENGES IN LLM EVALUATION

A. Reliability and Reproducibility

Several factors complicate reliable evaluation of LLMs:

Prompt Sensitivity: Minor variations in prompts can lead to significant performance differences, making evaluation results highly dependent on prompt engineering (Liu et al., 2023).

Stochasticity: Random sampling during generation introduces variability in outputs, complicating consistent evaluation (Zhao et al., 2023).

Benchmark Saturation: As models improve, existing benchmarks quickly become saturated, providing diminishing value for distinguishing between model capabilities (Kiela et al., 2021).

Replication Challenges: Closed-source models and limited documentation of evaluation methodologies make reproducing reported results difficult for independent researchers (Mitchell et al., 2022).

B. Evaluating Emergent Capabilities

Modern LLMs exhibit emergent capabilities that are challenging to evaluate systematically:

In-Context Learning: The ability to adapt to new tasks from examples provided in the prompt, without parameter updates (Brown et al., 2020). Evaluating this capability requires carefully designed few-shot prompts that may not generalize across models.

Chain-of-Thought Reasoning: The capacity to perform explicit multi-step logical reasoning (Wei et al., 2022). Current metrics often focus on final answers rather than the reasoning process, missing important aspects of model performance.

Tool Use and Agency: The ability to effectively use external tools, APIs, or engage in agentic behaviors involves complex interaction patterns that current benchmarks rarely capture (Schick et al., 2023).

C. Alignment Evaluation

Evaluating alignment with human values presents unique challenges:

Preference Misalignment: What models are optimized for may diverge from actual human preferences (Ouyang et al., 2022).

Safety-Performance Tradeoffs: Safety constraints may reduce performance on certain tasks, creating complex evaluation tradeoffs (Bai et al., 2022).

Cultural and Demographic Biases: Preferences used for alignment may reflect specific cultural or demographic biases, requiring more diverse evaluation approaches (Santurkar et al., 2023).

V. MULTI-DIMENSIONAL EVALUATION FRAMEWORKS

A. Beyond Task Performance

Comprehensive evaluation must consider multiple dimensions:

Safety Evaluation: Assessing propensity to generate harmful, biased, or misleading content:

- **Red-Teaming:** Systematic adversarial testing to identify model vulnerabilities (Dinan et al., 2022).
- **Toxicity Benchmarks:** Standardized datasets for measuring harmful outputs (Gehman et al., 2020).
- **Bias Evaluation:** Measuring disparate performance across demographic groups (Blodgett et al., 2020).

Efficiency Metrics:

- **Inference Speed:** Time required to generate responses, particularly important for real-time applications.
- **Memory Requirements:** Resources needed during inference, affecting deployment options.
- **Energy Consumption:** Environmental impact of model training and inference (Patterson et al., 2021).

Robustness Assessment:

- **Out-of-Distribution Performance:** Ability to handle inputs significantly different from training data (Hendrycks et al., 2020).
- **Adversarial Robustness:** Resistance to inputs designed to manipulate model behavior (Wang et al., 2023).

B. HELM: A Case Study in Holistic Evaluation

The Holistic Evaluation of Language Models (HELM) framework (Liang et al., 2022) represents a significant advancement in comprehensive LLM evaluation. HELM evaluates models across multiple scenarios and metrics:

Scenarios: 42 scenarios across 7 categories (question answering, information retrieval, summarization, toxicity, etc.)

Metrics: 7 metrics including accuracy, calibration, robustness, fairness, bias, efficiency, and toxicity.

Findings from HELM reveal important insights:

- Performance across tasks is not strongly correlated, with models showing different strengths.
- Larger models generally improve performance but with diminishing returns on certain tasks.
- Safety-performance trade-offs are evident, with more aligned models sometimes showing lower performance on certain tasks.

C. The Role of Human Evaluation

Despite advances in automated evaluation, human judgment remains crucial:

Expert Evaluation: Domain specialists assessing performance on specialized tasks:

- Medical professionals evaluating clinical reasoning (Singhal et al., 2022)
- Legal experts assessing legal analysis capabilities (Guha et al., 2023)

Crowdsourced Evaluation: Larger-scale evaluation using platforms like Mechanical Turk:

- Benefits: Larger sample sizes, diverse perspectives
- Limitations: Quality control challenges, potential biases (Clark et al., 2021)

Hybrid Approaches: Combining automated metrics with targeted human evaluation:

- Initial screening using automated metrics
- Focused human evaluation on specific capabilities or edge cases

VI. FUTURE DIRECTIONS AND RECOMMENDATIONS

A. Developing More Robust Evaluation Methodologies

Dynamic Benchmarks: Continuously evolving benchmarks that adapt as models improve, preventing benchmark saturation (Kiela et al., 2021).

Adversarially-Created Benchmarks: Datasets specifically designed to challenge current model capabilities, identifying blind spots and limitations (Wallace et al., 2022).

Process-Based Evaluation: Assessing not just final outputs but the reasoning processes and intermediate steps, particularly important for complex reasoning and problem-solving tasks (Wei et al., 2022).

B. Standardization Efforts

Evaluation Protocols: Developing standardized protocols for comprehensive evaluation across the industry:

- Common prompt formats and evaluation methodologies
- Shared benchmarks with clear versioning
- Transparent reporting requirements

Open Evaluation Platforms: Community-driven platforms for collaborative benchmarking:

- Democratizing access to evaluation resources
- Enabling independent verification of claimed results
- Supporting continuous model assessment

C. Interdisciplinary Approaches

Sociotechnical Evaluation: Assessing how models function within broader social systems:

- Deployment context consideration
- User interaction patterns
- Organizational impacts

Long-term Impact Assessment: Methods for evaluating longer-term consequences of LLM deployment:

- Labor market effects
- Educational impacts
- Information ecosystem influence

VII. CONCLUSION

This comprehensive review has examined the landscape of LLM performance evaluation, from traditional metrics to emerging multi-dimensional frameworks. The performance comparison across leading models reveals both significant advances in capability and persistent challenges in reliable evaluation.

Several key conclusions emerge:

1. **Evaluation Complexity:** As LLMs grow in capability, evaluation becomes increasingly multi-faceted, requiring assessment across dimensions of performance, safety, efficiency, and social impact.
2. **Model Differentiation:** While frontier models (GPT-4, Claude 3, Gemini) show comparable performance on many benchmarks, they exhibit distinctive strengths in specific areas like reasoning, coding, and truthfulness.

3. **Methodological Gaps:** Current evaluation approaches struggle to systematically assess emergent capabilities, alignment with human values, and performance in dynamic, interactive settings.
4. **Standardization Need:** The field would benefit from more standardized, transparent evaluation protocols that enable meaningful comparison across models and track progress over time.

Future research should focus on developing more dynamic, comprehensive evaluation frameworks that can keep pace with rapid advances in model capabilities. Additionally, greater emphasis on process-based evaluation, long-term impact assessment, and interdisciplinary approaches will be crucial for understanding these increasingly powerful and complex systems. As LLMs continue to evolve and find applications across sectors, robust evaluation methodologies will remain essential not only for technical advancement but also for responsible development and deployment of these transformative technologies.

REFERENCES

1. Bai, Y., Jones, A., Ndousse, K., Askell, A., Chen, A., DasSarma, N., ... & Irving, G. (2022). Training a Helpful and Harmless Assistant with Reinforcement Learning from Human Feedback. arXiv preprint arXiv:2204.05862.
2. Blodgett, S. L., Barocas, S., Daumé III, H., & Wallach, H. (2020). Language (Technology) is Power: A Critical Survey of "Bias" in NLP. Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics.
3. Brown, T. B., Mann, B., Ryder, N., Subbiah, M., Kaplan, J., Dhariwal, P., ... & Amodei, D. (2020). Language Models are Few-Shot Learners. Advances in Neural Information Processing Systems.
4. Callison-Burch, C., Osborne, M., & Koehn, P. (2006). Re-evaluating the Role of BLEU in Machine Translation Research. 11th Conference of the European Chapter of the Association for Computational Linguistics.
5. Chen, M., Tworek, J., Jun, H., Yuan, Q., Pinto, H. P. D. O., Kaplan, J., ... & Zaremba, W. (2021). Evaluating Large Language Models Trained on Code. arXiv preprint arXiv:2107.03374.
6. Clark, E., August, T., Serrano, S., Haduong, N., Gururangan, S., & Smith, N. A. (2021). All That's 'Human' Is Not Gold: Evaluating Human Evaluation of Generated Text. Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics.
7. Dinan, E., Abercrombie, G., Bergman, A. S., Spruit, S., Hupkes, D., Kiela, D., ... & Weidinger, L. (2022). SafetyKit: First Aid for Measuring Safety in Open-domain Conversational Systems. Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics.
8. Gao, L., Tow, J., Biderman, S., Black, S., DiPofi, A., Foster, C., ... & Leahy, C. (2023). A Framework for Rigorous Evaluation of Language Models. arXiv preprint arXiv:2306.05685.
9. Gehman, S., Gururangan, S., Sap, M., Choi, Y., & Smith, N. A. (2020). RealToxicityPrompts: Evaluating Neural Toxic Degeneration in Language Models. Findings of the Association for Computational Linguistics: EMNLP 2020.
10. Guha, A., Garimella, A., Jain, N., Miresghallah, F., Donnelly, L., Wang, B., ... & McKeown, K. (2023). LEGAL-BERT: LegalBench Evaluation for Reasoning and-Thinking. arXiv preprint arXiv:2308.11462.
11. Hendrycks, D., Basart, S., Mu, N., Kadavath, S., Wang, F., Dorundo, E., ... & Steinhardt, J. (2021). Measuring Massive Multitask Language Understanding. International Conference on Learning Representations.
12. Hendrycks, D., Liu, X., Wallace, E., Dziedzic, A., Krishnan, R., & Song, D. (2020). Pretrained Transformers Improve Out-of-Distribution Robustness. Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics.
13. Kiela, D., Firooz, H., Mohan, A., Goswami, V., Singh, A., Ringshia, P., & Testuggine, D. (2021). The Dynabench Platform for Dynamic Dataset Creation. Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics.
14. Liang, P., Bommasani, R., Lee, T., Tsipras, D., Soylu, D., Yasunaga, M., ... & Hashimoto, T. (2022). Holistic Evaluation of Language Models. arXiv preprint arXiv:2211.09110.
15. Liu, P., Yuan, W., Fu, J., Jiang, Z., Hayashi, H., & Neubig, G. (2023). Pre-train, Prompt, and Predict: A Systematic Survey of Prompting Methods in Natural Language Processing. ACM Computing Surveys.
16. Mitchell, M., Wu, S., Zaldívar, A., Barnes, P., Vasserman, L., Hutchinson, B., ... & Gebru, T. (2022). Model Cards for Model Reporting. Proceedings of the Conference on Fairness, Accountability, and Transparency.
17. Nori, H., King, N., McKinney, S. M., Carignan, D., & Horvitz, E. (2023). Capabilities of GPT-4 on Medical Challenge Problems. arXiv preprint arXiv:2303.13375.
18. Ouyang, L., Wu, J., Jiang, X., Almeida, D., Wainwright, C., Mishkin, P., ... & Lowe, R. (2022). Training Language Models to Follow Instructions with Human Feedback. Advances in Neural Information Processing Systems.
19. Patterson, D., Gonzalez, J., Le, Q., Liang, C., Munguia, L. M., Rothchild, D., ... & Dean, J. (2021). Carbon Emissions and Large Neural Network Training. arXiv preprint arXiv:2104.10350.
20. Radford, A., Wu, J., Child, R., Luan, D., Amodei, D., & Sutskever, I. (2019). Language Models are Unsupervised Multitask Learners. OpenAI Blog.
21. Santurkar, S., Tsipras, D., Sohoni, N., Vogel, B., Panda, S., Sridhar, K., ... & Hashimoto, T. (2023). Whose Opinions Do Language Models Reflect? arXiv preprint arXiv:2303.17548.

22. Schick, T., Dwivedi-Yu, J., Dessì, R., Raileanu, R., Lomeli, M., Zettlemoyer, L., ... & Scialom, T. (2023). Toolformer: Language Models Can Teach Themselves to Use Tools. arXiv preprint arXiv:2302.04761.
23. Singhal, K., Azizi, S., Tu, T., Mahdavi, S. S., Wei, J., Chung, H. W., ... & Shlens, J. (2022). Large Language Models Encode Clinical Knowledge. arXiv preprint arXiv:2212.13138.
24. Srivastava, A., Rastogi, A., Rao, A., Shoeb, A. A. M., Abid, A., Fisch, A., ... & Zoph, B. (2022). Beyond the Imitation Game: Quantifying and Extrapolating the Capabilities of Language Models. arXiv preprint arXiv:2206.04615.
25. Wallace, E., Feng, S., Kandpal, N., Gardner, M., & Singh, S. (2022). Universal Adversarial Triggers for Attacking and Analyzing NLP. Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing.
26. Wang, A., Pruksachatkun, Y., Nangia, N., Singh, A., Michael, J., Hill, F., ... & Bowman, S. R. (2019). SuperGLUE: A Stickier Benchmark for General-Purpose Language Understanding Systems. Advances in Neural Information Processing Systems.
27. Wang, A., Singh, A., Michael, J., Hill, F., Levy, O., & Bowman, S. R. (2018). GLUE: A Multi-Task Benchmark and Analysis Platform for Natural Language Understanding. Proceedings of the 2018 EMNLP Workshop BlackboxNLP.
28. Wang, W., Wang, J., Shen, A., Wang, Y., Li, Y., Zhang, T., ... & Li, H. (2023). DecodingTrust: A Comprehensive Assessment of Trustworthiness in GPT Models. arXiv preprint arXiv:2306.11698.
29. Wei, J., Bosma, M., Zhao, V. Y., Guu, K., Yu, A. W., Lester, B., ... & Le, Q. V. (2022). Chain-of-Thought Prompting Elicits Reasoning in Large Language Models. Advances in Neural Information Processing Systems.
30. Zhao, W. X., Zhou, K., Li, J., Tang, T., Wang, X., Hou, Y., ... & Wen, J. R. (2023). A Survey of Large Language Models. arXiv preprint arXiv:2303.18223.