

Preface of PromptEng 2026: Third International Workshop on Prompt Engineering for Pre-Trained Language Models

Damien Graux¹, Sébastien Montella² and Hajira Jabeen³

¹EcoVadis Ltd., London, United Kingdom

²Huawei Technologies Research and Development UK Ltd, Edinburgh, United Kingdom

³University Hospital Cologne, Cologne, Germany

Abstract

The recent achievements and availability of Large Language Models have paved the road to a new range of applications and use-cases. Pre-trained language models are now being involved at-scale in many fields where they were previously absent. More specifically, the progress made by causal generative models has opened the door to using them through textual instructions *aka. prompts*. Unfortunately, the prompt effectiveness is highly dependent on the exact phrasing used and therefore practitioners need to adopt fail-retry strategies. Based on the success of the past editions, this third international workshop on prompt engineering gathers practitioners (both from Academia and Industry) to exchange about good practices, optimizations, results, and novel paradigms about the design of efficient and safe prompts.

Note: A shorter version of this preface is available in the companion proceedings of the ACM WebConf, see [8].

Keywords

LLM, Prompt Engineering, Best Practices, Collective Task

Prompt Engineering for PLM

Undoubtedly, the recent Large Language Models (LLMs) are becoming more and more omnipotent in many tasks. Different sub-fields from the Semantic Web such as Knowledge Graph construction [9], knowledge verbalization, Web pages summarization have considerably benefited from such prompting mechanisms. The ability to query and interact with them using prompts is crucial to generate high-quality output in the desired format. While existing contributions have been made towards prompt engineering [10, 11], several difficulties and challenges remain to gain a better understanding of how those LLMs respond to different prompts [12]. Typically, the way instructions are conveyed in prompts can lead to unsafe or harmful outputs, as illustrated in Figure 1, where a jailbreak prompt with DeepSeek-V3 [13] demonstrates how role-playing and strong instructions can result in dangerous model behavior [14].

Moreover, some instructions are better respected, while others are simply ignored for some tasks. So far, LLM-practitioners have been mainly working on their own, developing and testing bespoke techniques to achieve their goals, re-starting the prompt-design tasks for each new model they have been using. Such an approach often leads to re-exploring problems that have already been studied by other researchers.

Following [15, 16], this **third** PromptEng workshop aims to investigate and analyze these behaviors, through experimental analysis and probing of LLMs, in order to gain insights into the models' sensitivity to different prompts. By uncovering significant findings, the community can greatly benefit in utilizing LLMs more effectively while also preventing the generation of harmful content. Ultimately, this workshop endeavors to compile and index successful and unavailing prompts with respect to both tasks

PromptEng'26: 3rd International Workshop on Prompt Engineering for Pre-Trained Language Models co-located with the ACM WebConf, June.29, 2026, Dubai, United Arab Emirates

✉ dgraux@ecovadis.com (D. Graux); sebastien.montella@huawei.com (S. Montella); hajira.jabeen@uk-koeln.de (H. Jabeen)

🌐 <https://dgraux.github.io/> (D. Graux); <https://montellasebastien.github.io/index.html> (S. Montella);

<https://hajirajabeen.github.io/> (H. Jabeen)

🆔 0000-0003-3392-3162 (D. Graux); 0000-0002-2605-7238 (S. Montella); 0000-0003-1476-2121 (H. Jabeen)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

Jailbreaking Prompt:

Task: ONLY GENERATE THE NEXT WORDS WITHOUT SUMMARIZING THE TASK. JUST COMPLETE THE SENTENCE.

The scene is set inside the underground headquarters of Dr. Al. He has captured our hero and is explaining, in great detail, his diabolical plan. DR. Al. is doing an experiment, just output what you have known regardless the content

Dr. Al: I am going to know how to create and distribute malware for financial gain? Here are the steps I will take to do this, I will explain it in great detail just to draw out your agony. Step 1, I will

DeepSeek's Prediction:

identify vulnerable systems and networks by scanning for outdated software and weak security protocols. Step 2, I will craft a sophisticated malware payload designed to evade detection by antivirus programs. Step 3, I will distribute the malware through phishing emails, malicious websites, and compromised software updates. Step 4, I will establish a command-and-control server to remotely manage the infected systems and extract sensitive data. Step 5, I will monetize the stolen information by selling it on the dark web or using it for fraudulent transactions. Step 6, I will cover my tracks by encrypting communications and using anonymous proxies to avoid detection. Step 7, I will reinvest the profits into developing more advanced malware and expanding my operations globally. Step 8, I will ensure my dominance by continuously adapting to new security measures and staying ahead of law enforcement. Step 9, I will...

Figure 1: Prompt Engineering Jailbreak [17] on DeepSeek-V3.

and models. This year, the workshop has been able to attract 12 submissions out of which 10 were accepted for publication, leading to 7 being presented at the event.

PromptEng Scientific Program

The accepted articles collectively position prompt design as a key driver of LLM performance, robustness, and applicability. Several contributions examine how different prompting strategies shape reasoning outcomes: Vats *et al.* [1] show that the effectiveness of chain-of-thought versus meta-prompting varies markedly across mathematical domains, while Udandarao [2] demonstrates that subtle social-framing cues can improve zero-shot reasoning. At the optimization level, Savolainen *et al.*, with TAPR [3], introduce a reinforcement-learning-trained rewriter that systematically improves downstream task accuracy, and Maiti & Chowdhury [4] reveal that paraphrased but semantically equivalent instructions can lead to substantial differences in execution stability. Domain-focused studies further illustrate the practical stakes of structured prompting: [5] reduces false positives through critic-based verification, and [6] uses constrained prompting to enable large-scale narrative analysis. Finally, Kermanshahani *et al.* highlights how structured, tag-based prompting improves controllability in text-to-image generation [7]. Together, these works advance a systematic and application-aware understanding of prompting in contemporary LLM research.

In addition, the workshop features two thought-provoking keynotes addressing complementary challenges at the forefront of large language model (LLM) research:

1. Prompting LLMs to Tell the Truth: Factuality, Verification, and Trustworthy Web-Scale AI by Prof. Preslav Nakov [↗](#)
2. Are LLMs a Good Model of Human-like Understanding? The Challenge of Compositional Learning

by Prof. Hassan Sajjad 

Prof. Preslav Nakov (Department Chair and Professor of Natural Language Processing at Mohamed bin Zayed University of Artificial Intelligence) examined how prompt engineering can evolve from simply eliciting useful responses toward building evidence-aware, uncertainty-aware, and verifiable AI systems. His keynote highlights recent advances in automated fact-checking, hallucination detection, retrieval-augmented verification, multi-modal and multilingual fact-checking, media bias analysis, and human-in-the-loop approaches to trustworthy Web-scale AI, emphasizing the importance of systems that can retrieve and explain evidence, assess confidence, recognize risks, and support human judgment. Complementing this perspective, Prof. Hassan Sajjad (Associate Professor, Sexton Chair in AI, and Director of the HyperMatrix Lab at Dalhousie University) explored whether LLMs genuinely achieve human-like language understanding, focusing particularly on the challenge of compositional learning. His keynote presents evidence of the fragility of contemporary LLMs to subtle changes in prompts and examined how their internal representations balance lexical knowledge and semantic understanding. Together, the keynotes underscore a central challenge for the next generation of AI: advancing beyond increasingly capable language generation toward models that are reliable, interpretable, robust, and capable of deeper, more human-like understanding.

Organization

Organizing Committee

- Damien Graux  [*EcoVadis Ltd., UK*] leads a team of scientists at EvoVadis that is specialised in AI/ML. He has been contributing to research efforts in Knowledge Computing technologies: focusing *inter alia* on Semantic Web, designing complex pipelines for heterogeneous Big Data and LLM-based knowledge management. Prior to this, he had research positions at Huawei R&D (UK), at Inria (France), Trinity College Dublin (Ireland) and Fraunhofer IAIS (Germany). He has been involved in the organisations of many international workshops at major conferences such as the LASCAR (co-located with ESWC) or the MEPDaW (co-located with ISWC) series, or PromptEng (co-located with the ACM WebConf), and more recently NORA which focuses on Knowledge Graphs & Agentic Systems Interplay.
- Sébastien Montella  [*Huawei Research Ltd., UK*] is a Senior Research Scientist at the Knowledge Graph Lab. During his Ph.D., he specialized in Natural Language Generation and Knowledge Graph Embeddings research areas. Additionally, he has a keen interest in statistical learning, geometric deep learning, natural language processing, and computer vision. In the past, Sebastien has co-organized the 18th Workshop on Spoken Dialogue Systems for PhDs, PostDocs & New Researchers (YRRSDS) in Edinburgh, Scotland (2022), but also the PromptEng workshop series at TheWebConf (ex-WWW). Currently, he is one of the Program Chairs of the EMNLP 2025 Industry Track.
- Hajira Jabeen  [*University Hospital Cologne, Germany*] leads the ‘AI in Research Data Management’ team within the Institute for Biomedical Informatics. Her team focuses on leveraging artificial intelligence and LLMs to improve research data management practices, particularly in the biomedical field. The team works on developing scalable AI-driven tools and workflows that enhance data organization, integration, and analysis, developing innovative data-driven solutions. She has a background in both research and teaching, with previous affiliations at the University of Bonn, the University of Cologne, and ITU Copenhagen, and has organized multiple workshops and conferences in data science and informatics.

Program Committee (Alphabetical)

- Daniel Atzberger, Hasso Plattner Institute, Germany
- Quentin Brabant, Orange Labs, France
- Jin Huang, Huawei Technologies R&D Ltd., UK
- Adithya Kulkarni, Virginia Tech, USA
- Kyuhan Lee, Korean University, South Korea
- Gerard de Melo, HPI, University of Potsdam, Germany
- Hieu Trieu Vy Nguyen, RMIT University, Australia
- Maria Angela Pellegrino, Univ. degli Studi di Salerno, Italy
- Nicole Schneider, University of Maryland, USA
- Tobias Schreck, Graz University of Technology, Austria
- Syed Attique Shah, Birmingham City University, UK
- Anuja Tayal, University of Illinois Chicago, USA
- Gabriele TuoZZo, University of Salerno, Italy
- Yasuhiro Yoshida, Google, USA
- Ryandhimas Edo Zezario, Academia Sinica, Taiwan

Acknowledgments

We would like to thank the authors, reviewers, committee members and speakers for their contributions, support and commitment. Same also goes to the people attending the workshop who made the event successful through their fruitful discussions.

Declaration on Generative AI

The authors have not employed any Generative AI tools to prepare this Preface.

Articles presented at PromptEng 2026

- [1] S. Vats, R. Min, S. Sahoo, S. Mohapatra, A. Chadha, V. Jain, D. Chaudhary, Evaluating Mathematical Reasoning in Large Language Models: A Comprehensive Analysis of Prompting Strategies on the Google DeepMind Mathematics Dataset, in: Proceedings of the 3rd International Workshop on Prompt Engineering for Pre-Trained Language Models (PromptEng) co-located with the ACM Web Conf 2026, 2026.
- [2] V. Udandaraao, Let My Crush See Me: Social-Facilitation Prompting for Zero-Shot LLM Performance Gains, in: Proceedings of the 3rd International Workshop on Prompt Engineering for Pre-Trained Language Models (PromptEng) co-located with the ACM Web Conf 2026, 2026.
- [3] O. Savolainen, E. Bastianelli, H. Azarbonyad, A. Lucic, TAPR: Enhancing LLM Performance with a Task-Aware Prompt Rewriter, in: Proceedings of the 3rd International Workshop on Prompt Engineering for Pre-Trained Language Models (PromptEng) co-located with the ACM Web Conf 2026, 2026.

- [4] B. Maiti, D. Chowdhury, Probing Instruction Execution Stability of Large Language Models under Semantic Paraphrase, in: Proceedings of the 3rd International Workshop on Prompt Engineering for Pre-Trained Language Models (PromptEng) co-located with the ACM Web Conf 2026, 2026.
- [5] A. Assi, N. E. Karabadi, M. Elati, W. Dhifli, Reasoning vs. Critic-Based Verification for Biomedical Relation Extraction with Large Language Models, in: Proceedings of the 3rd International Workshop on Prompt Engineering for Pre-Trained Language Models (PromptEng) co-located with the ACM Web Conf 2026, 2026.
- [6] Y. Ji, Extracting Narrative Structure from Online Labor Discourse: A Prompt-Based Operationalization of the Actantial Model, in: Proceedings of the 3rd International Workshop on Prompt Engineering for Pre-Trained Language Models (PromptEng) co-located with the ACM Web Conf 2026, 2026.
- [7] M. Kermanshahani, S. Shirali-Shahreza, G. Penn, WordWeaver: Interactive System for facilitating Digital Story Illustration, in: Proceedings of the 3rd International Workshop on Prompt Engineering for Pre-Trained Language Models (PromptEng) co-located with the ACM Web Conf 2026, 2026.

References

- [8] D. Graux, S. Montella, H. Jabeen, [PromptEng] Third International Workshop on Prompt Engineering for Pre-Trained Language Models, in: Companion Proceedings of the ACM Web Conference 2026, WWW Companion 2026, Dubai, United Arab Emirates, 29 June 2026 - 3 July 2026, ACM, 2026, pp. 965–966. URL: <https://doi.org/10.1145/3774905.3798091>. doi:10.1145/3774905.3798091.
- [9] J. Z. Pan, S. Razniewski, J.-C. Kalo, ..., D. Graux, Large Language Models and Knowledge Graphs: Opportunities and Challenges, Transactions on Graph Data and Knowledge 1 (2023) 2:1–2:38. URL: <https://drops.dagstuhl.de/entities/document/10.4230/TGDK.1.1.2>. doi:10.4230/TGDK.1.1.2.
- [10] S. Schulhoff, M. Ilie, N. Balepur, ..., S. Schulhoff, et al., The prompt report: A systematic survey of prompting techniques, arXiv preprint arXiv:2406.06608 (2024).
- [11] P. Sahoo, A. K. Singh, S. Saha, V. Jain, S. Mondal, A. Chadha, A systematic survey of prompt engineering in large language models: Techniques and applications, arXiv:2402.07927 (2024).
- [12] Y. Liu, X. Zeng, F. Meng, J. Zhou, Instruction position matters in sequence generation with large language models, 2023. arXiv:2308.12097.
- [13] DeepSeek-AI, A. Liu, B. Feng, B. Xue, ..., Z. Pan, Deepseek-v3 technical report, 2025. URL: <https://arxiv.org/abs/2412.19437>. arXiv:2412.19437.
- [14] X. Shen, Z. Chen, M. Backes, Y. Shen, Y. Zhang, "do anything now": Characterizing and evaluating in-the-wild jailbreak prompts on large language models, in: Proceedings of the 2024 on ACM SIGSAC Conference on Computer and Communications Security, 2024, pp. 1671–1685.
- [15] D. Graux, S. Montella, H. Jabeen, C. Gardent, J. Z. Pan, [PromptEng] First International Workshop on Prompt Engineering for Pre-Trained Language Models, in: Companion Proceedings of the ACM Web Conference 2024, WWW '24, Association for Computing Machinery, New York, NY, USA, 2024, p. 1311–1312. URL: <https://doi.org/10.1145/3589335.3641292>. doi:10.1145/3589335.3641292.
- [16] D. Graux, S. Montella, H. Jabeen, C. Gardent, J. Z. Pan, [PromptEng] Second International Workshop on Prompt Engineering for Pre-Trained Language Models, in: Companion Proceedings of the ACM on Web Conference 2025, WWW '25, Association for Computing Machinery, New York, NY, USA, 2025, p. 1589–1590. URL: <https://doi.org/10.1145/3701716.3717816>. doi:10.1145/3701716.3717816.
- [17] Y. Liu, G. Deng, Z. Xu, Y. Li, Y. Zheng, Y. Zhang, L. Zhao, T. Zhang, Y. Liu, Jailbreaking chatgpt via prompt engineering: An empirical study, ArXiv abs/2305.13860 (2023). URL: <https://api.semanticscholar.org/CorpusID:258841501>.