

Benchmarking Generative AI for Requirements Management

Ruslan Bernijazov
& Nikhil Jha



About us

Speakers of the session



RUSLAN BERNIJAZOV

MD & Co-founder, AI Marketplace GmbH
10+ Years Experience in Systems Engineering
Fraunhofer IEM | Heinz Nixdorf Institute



NIKHIL JHA

AI Technology Generalist, AI Marketplace GmbH
Lead, Google Developers Group Paderborn
Fraunhofer IEM | Wipro



About us

Engineering slows down as complexity grows



The daily routine of an engineer:



The daily problems of an engineer:

- 1 Change impact analysis across requirements & architecture takes weeks
- 2 Engineering data is fragmented across tools
- 3 Critical artifacts (requirements, architectures, simulation models) are not AI-ready

About us

We're building the platform that operates agents in engineering



Current way

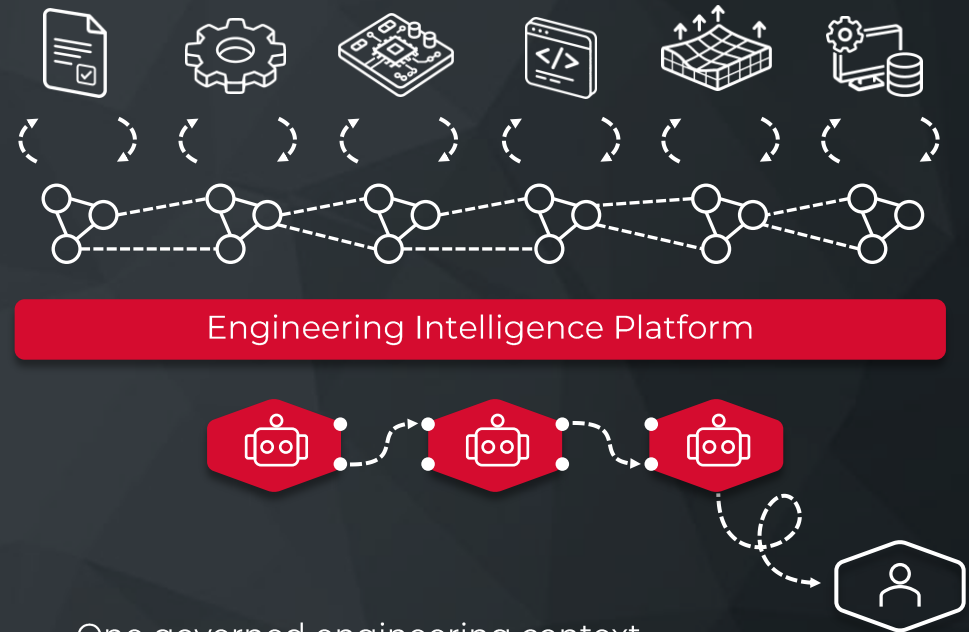
Requirements from customer Geometries from CAD Layouts from E-CAD Code from Git Results from Simulation Master data from the PLM



- Engineer searches ReqIF, SysML, PDFs
- Manual cross-tool reasoning
- High risk of missing impacts

Our vision

and why it's better



- One governed engineering context
- Agent analyzes impact
- Human approves → traceable decision

Challenges in Requirements Management

Manual Process

Error Prone

Regulations

Data Constraints

Compliance



Challenges in Requirements Management

Manual Process

Error Prone

Regulations

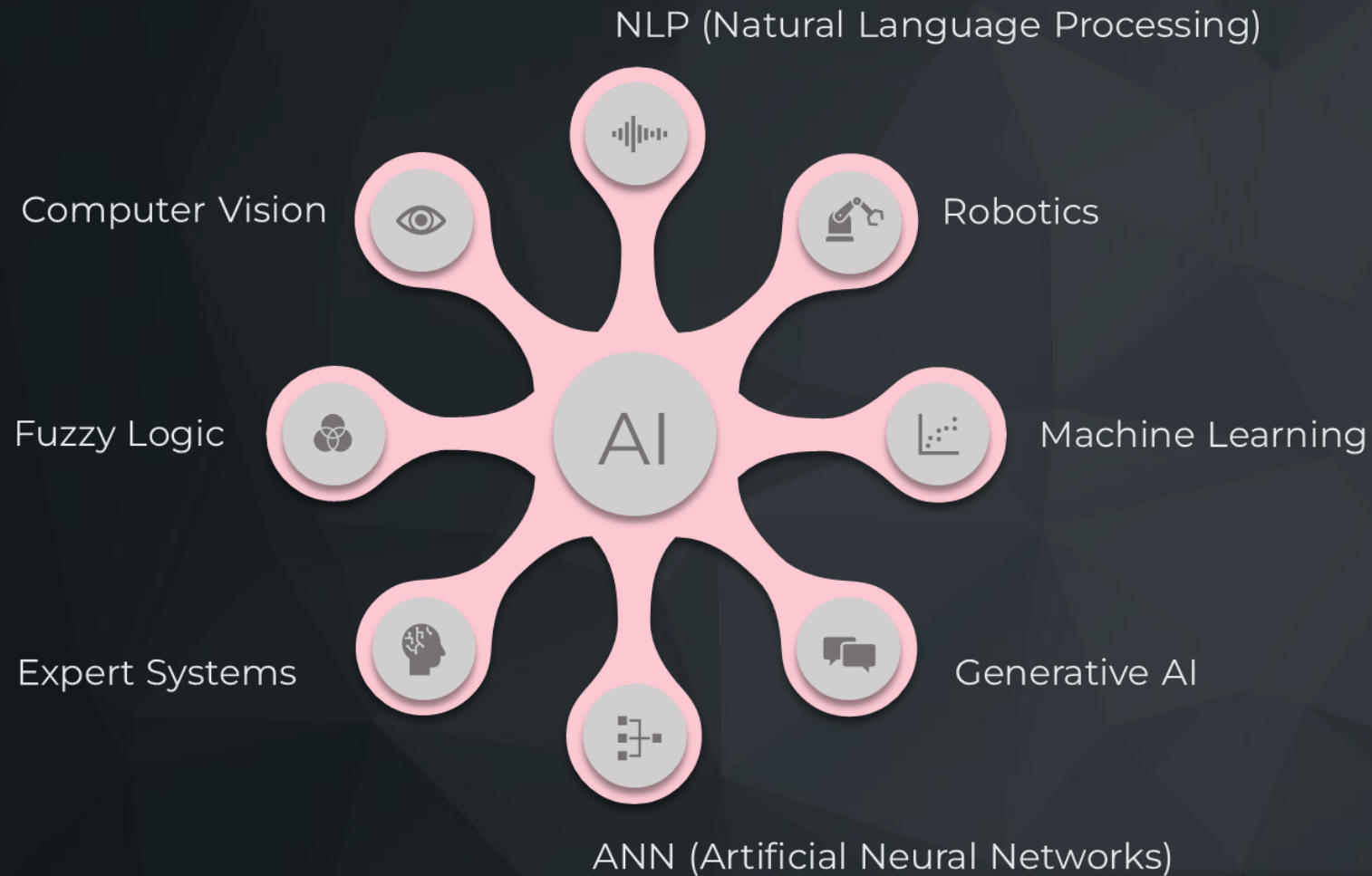
Data Constraints

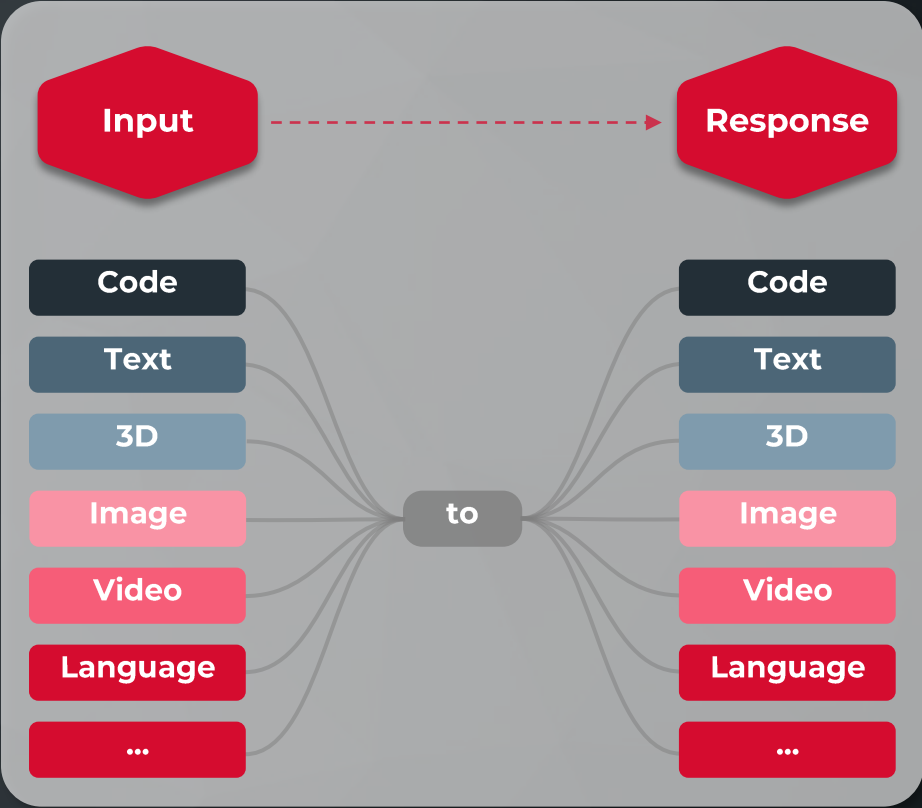
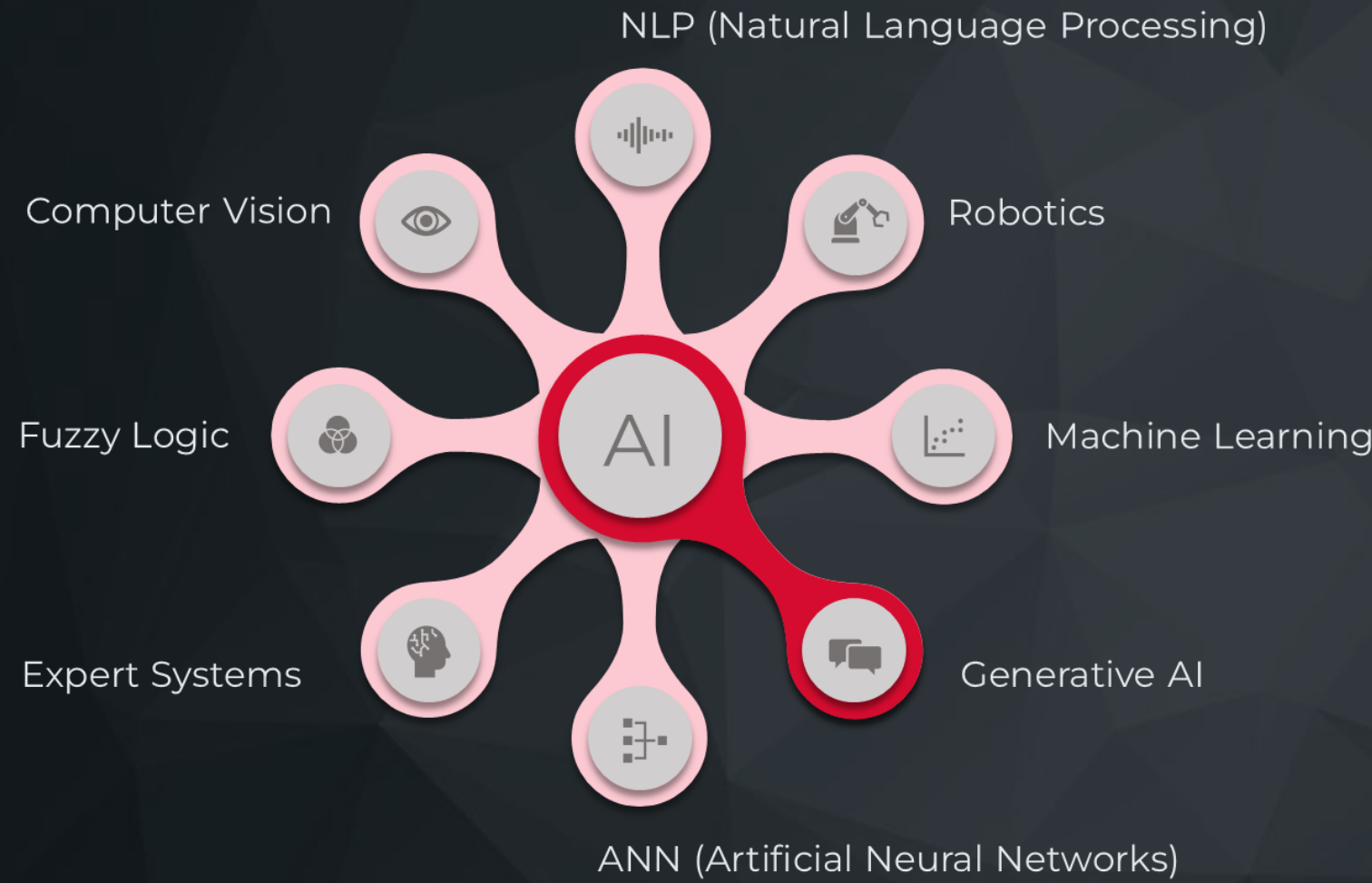
Compliance

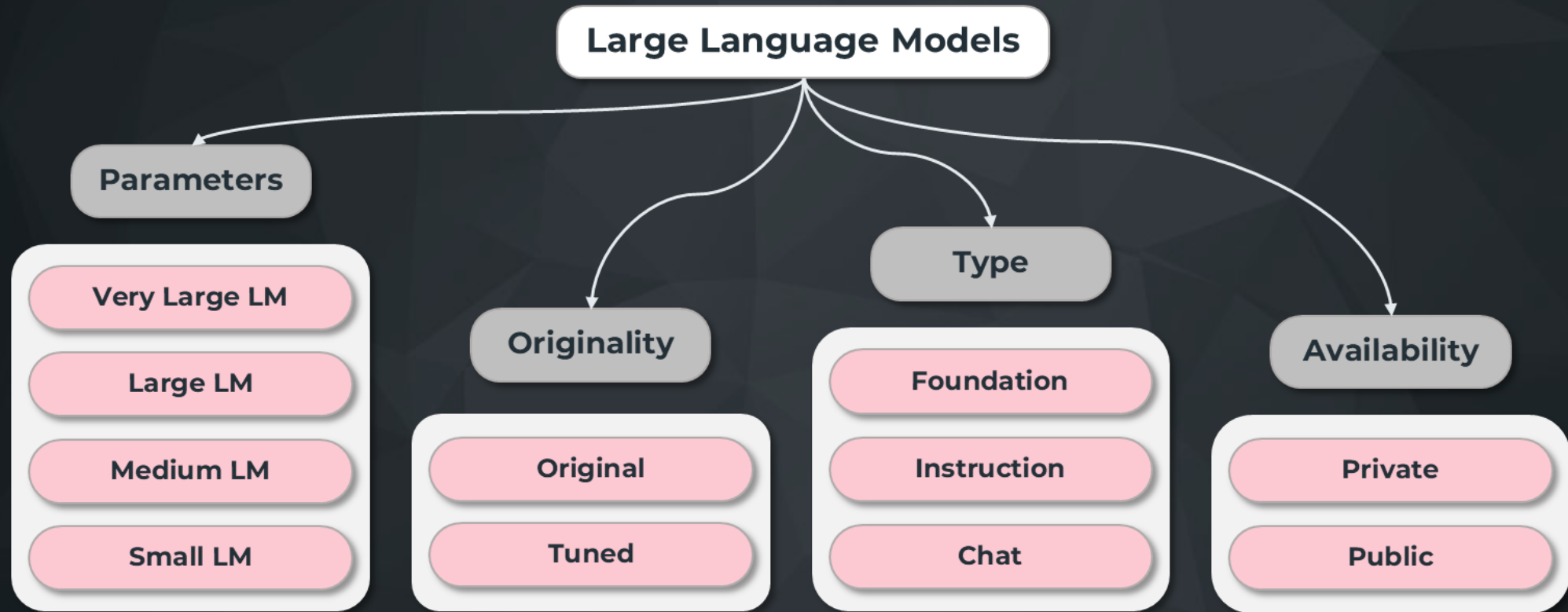


Artificial
Intelligence

Fields of Artificial Intelligence (AI)







Prompting Techniques



Zero-Shot
Prompting



Few-Shot
Prompting



Chain-of-
thought



Tree-of-
thought



Iterative
refinement



Feedback
loop



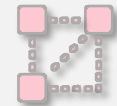
Prompt
chaining



Role-
playing



Maieutic
Prompting



Complexity
based
Prompting

Motivation

Why Generative AI ?



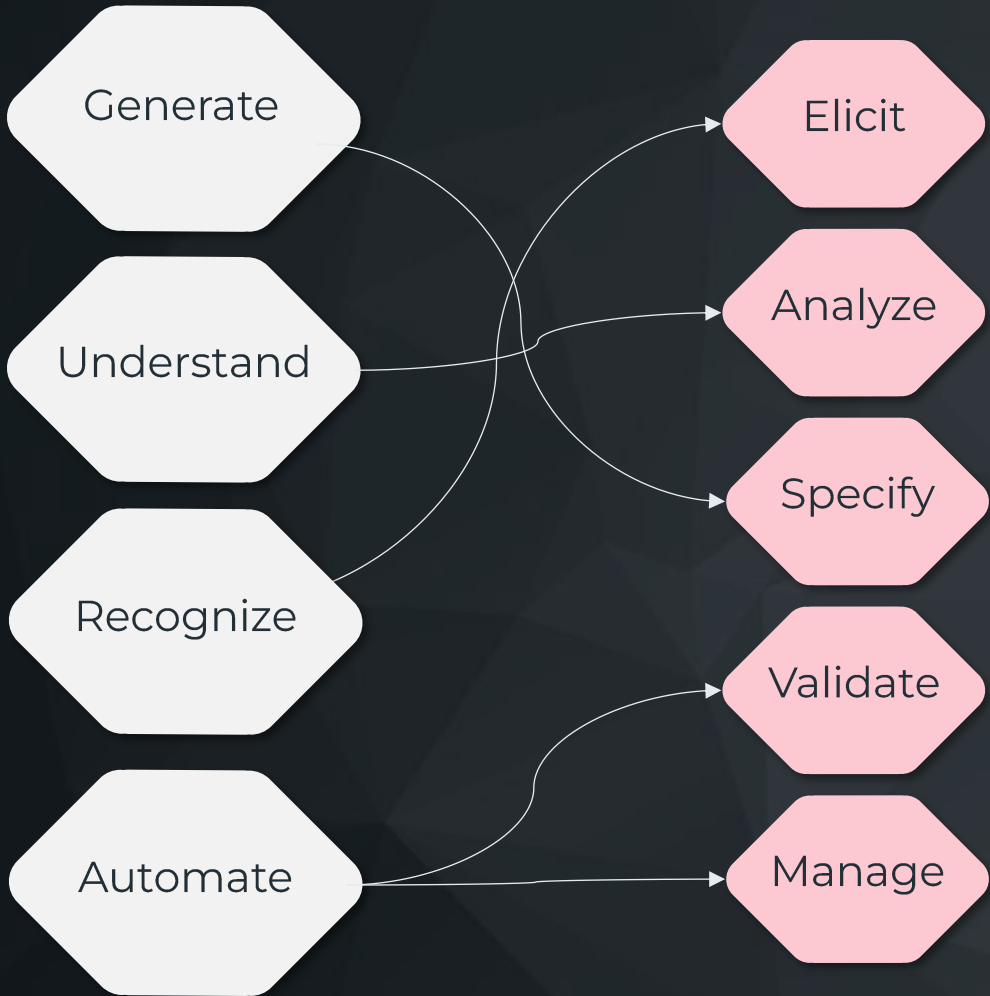
Generate

Understand

Recognize

Automate

Why Generative AI in RM?



Advancing Requirements Engineering Through Generative AI: Assessing the Role of LLMs

Chetan Arora, John Grundy, and Mohamed Abdelrazek

Abstract Requirements Engineering (RE) is a critical component including the elicitation, analysis, specification, and validation of requirements. Despite the importance of RE, it remains the complexities of communication, uncertainty in automation support. In recent years, large language models (LLMs) have shown significant promise in diverse domains, including code generation, and program understanding. This paper explores the role of LLMs in driving RE processes, aiming to improve requirements-related tasks. We propose key directions for research and development in using LLMs for RE requirements elicitation, analysis, specification, and validation, based on the results from a preliminary evaluation, in this context.

Keywords Requirements engineering · Generative AI (LLMs) · Natural language processing · Software development

1 Introduction

Requirements Engineering (RE) is arguably the most critical phase in the software development process, where the needs and constraints are analyzed, and documented to create a well-defined set of requirements. However, project teams often overlook or defer RE and its impact on project success [24]. The lack of effort and resources spent in RE includes:

© The Author(s), under exclusive license to Springer Nature A. Nguyen-Duc et al. (eds.), *Generative AI for Effective Software Development*, https://doi.org/10.1007/978-3-031-55642-5_6

Requirements Engineering Using Generative AI: Prompts and Prompting Patterns

Krishna Ronanki, Beatriz Cabrero-Daniel, Jennifer Horkoff, and Christian Berger

Abstract [Context] Companies are increasingly recognizing the value of automating Requirements Engineering (RE) tasks due to their resource nature. The advent of GenAI has made these tasks more amenable to automation thanks to its ability to understand and interpret context effectively. However, in the context of GenAI, prompt engineering is a critical factor. Despite this, we currently lack tools and methods to systematically determine the most effective prompt patterns to employ for a particular RE task. [Method] Two tasks related to requirements, specifically requirement elicitation and tracing, were automated using the GPT-3.5 turbo API. The evaluation involved assessing various prompts created using 5 patterns and implemented programmatically to perform the selected RE tasks on metrics such as precision, recall, accuracy, and F-Score. [Results] The results evaluate the effectiveness of the 5 prompt patterns' ability to make GPT-3.5 perform the selected RE tasks and offers recommendations on which pattern to use for a specific RE task. Additionally, it also provides an evaluation as a reference for researchers and practitioners who want to evaluate prompt patterns for different RE tasks.

Keywords Requirements engineering · Generative AI · Prompt pattern · Prompt engineering · Large language models

1 Introduction

Researchers in Requirements Engineering (RE) have been exploring the use of machine learning (ML) and deep learning (DL) methods for various RE tasks, including requirements classification, prioritization, tracing, ambiguity detection, and modelling [1]. However, the majority of existing ML/DL approaches are based on supervised learning, which requires a large amount of labeled data. This is often difficult to obtain in the RE domain, where requirements are often unstructured and the data is often noisy. In this paper, we propose a novel approach for RE tasks using Generative AI (GenAI). GenAI models, such as GPT-3, can generate text that is similar to human-generated text. This makes them a promising candidate for RE tasks, where the goal is to generate requirements that are consistent with the context. We evaluate the performance of GenAI models on RE tasks, and compare the results with traditional ML/DL approaches. Our results show that GenAI models can outperform traditional ML/DL approaches in RE tasks, particularly in the areas of requirements generation and tracing. This suggests that GenAI has the potential to revolutionize RE processes and outcomes. The integration of GenAI into RE presents both promising opportunities and significant challenges that necessitate systematic analysis and evaluation.

Objective: This paper presents a comprehensive systematic literature review (SLR) analyzing state-of-the-art applications and innovative proposals leveraging GenAI in RE. It surveys studies focusing on the utilization of GenAI to enhance RE processes while identifying key challenges and opportunities in this rapidly evolving field.

Method: A rigorous SLR methodology was used to conduct an in-depth analysis of 27 carefully selected primary studies. The review examined research questions pertaining to the application of GenAI across various RE phases, the models and techniques used, and the challenges encountered in implementation and adoption.

Results: The most salient findings include i) a predominant focus on the early stages of RE, particularly the elicitation and analysis of requirements, indicating potential for expansion into later phases; ii) the dominance of large language models, especially the GPT series, highlighting the need for diverse AI approaches; and iii) persistent challenges in domain-specific applications and the interpretability of AI-generated outputs, underscoring areas requiring further research and development.

Conclusions: The results highlight the critical need for comprehensive evaluation frameworks, improved human-AI collaboration models, and thorough consideration of ethical implications in GenAI-assisted RE. Future research should prioritize extending GenAI applications across the entire RE lifecycle, enhancing domain-specific capabilities, and developing strategies for responsible AI integration in RE practices.

Keywords: Generative AI, Requirements Engineering, Systematic Literature Review

Preprint submitted to Information and Software Technology September 12, 2024

arXiv:2409.06741v1 [cs.SE] 10 Sep 2024

Limitations of Generative AI



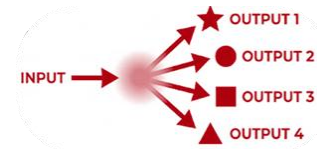
Hallucination



Coherence Loss



Black Box

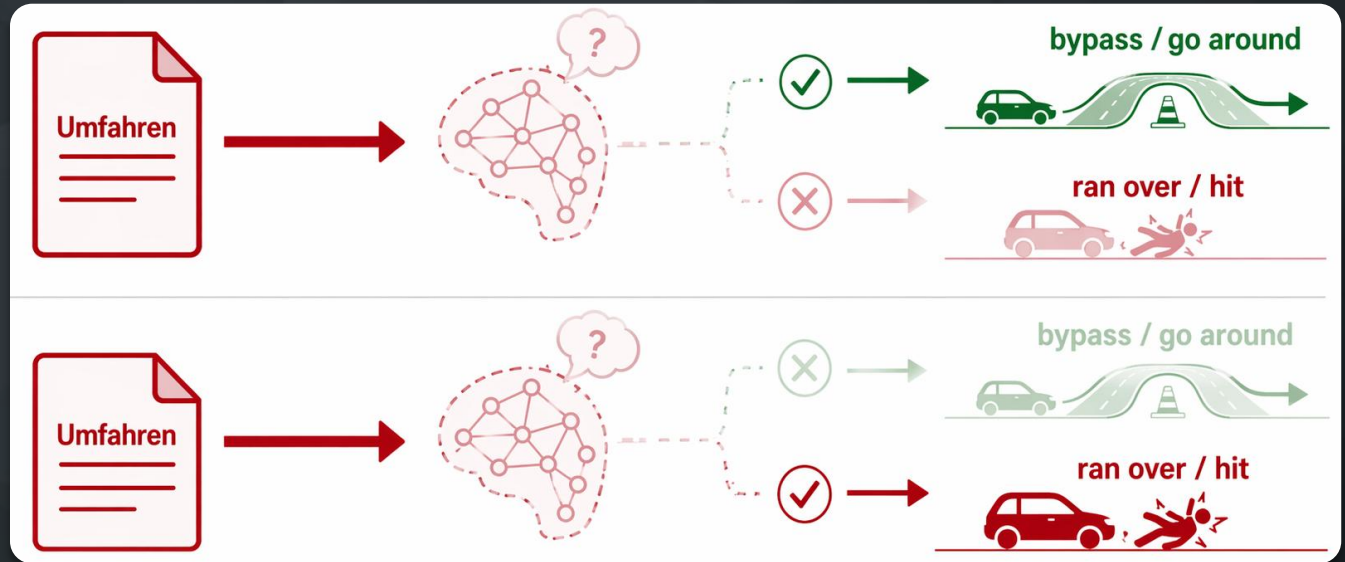


Irreproducibility

Limitations of Generative AI



Umfahren



THE GREAT GenAI RACE



BENCHMARKS

What is Benchmark ?

A scientific way for **evaluating and comparing** system performance **under standardized conditions**.



Workload

Methodology

Metrics



SWE-bench

Issue

data leak in GBDT due to warm start (This is about the non-histogram-based version of...

Codebase

- sklearn/ reqs.txt
- examples/ setup.cfg
- README.rst setup.py

Language Model

Generated PR

+20 -12

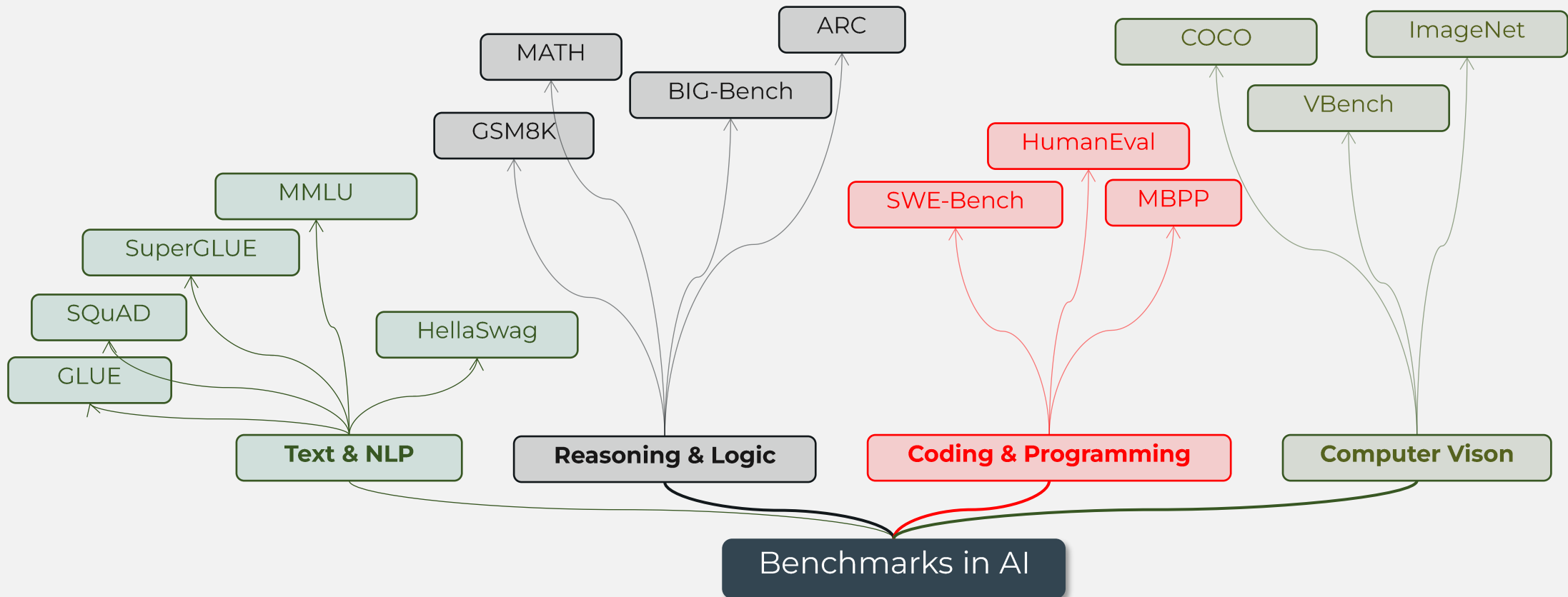
- sklearn
- gradient_boosting.py
- helper.py
- utils

Unit Tests

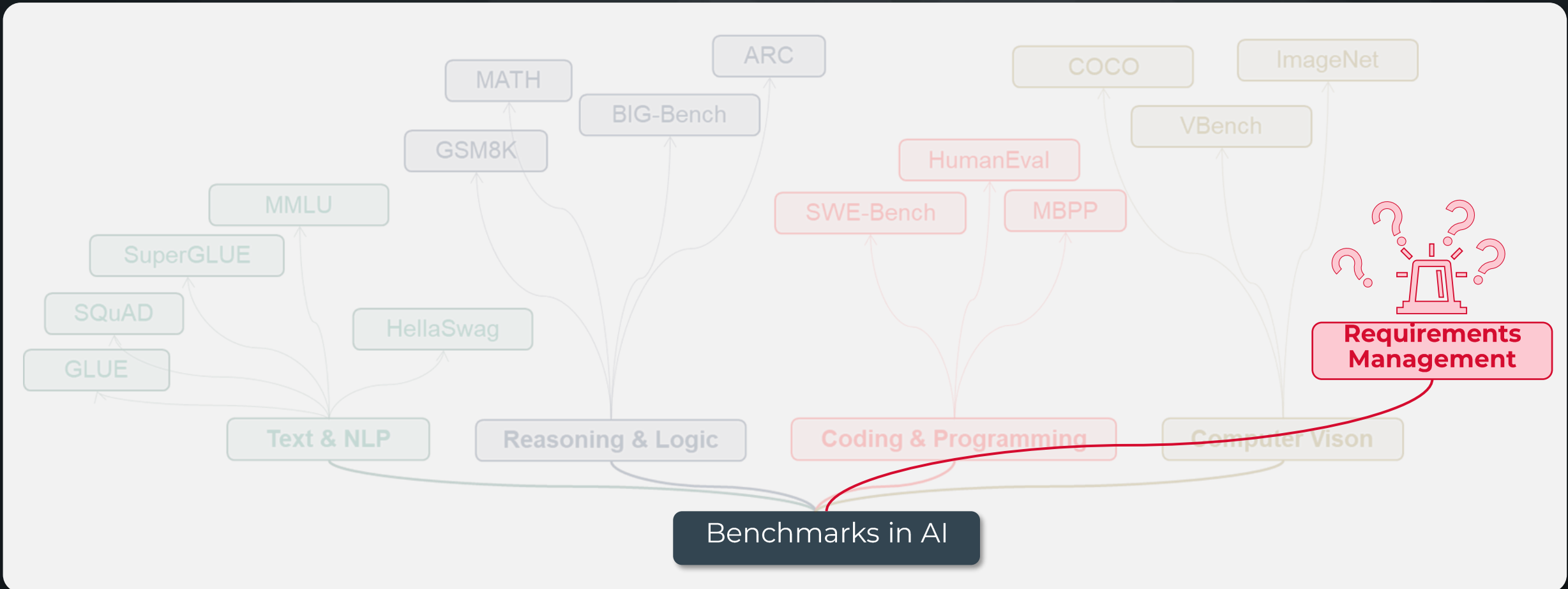
Pre PR	Post PR	Tests
✗	✓	join_struct_col
✗	✓	vstack_struct_col
✗	✓	dstack_struct_col
✓	✓	matrix_transform
✓	✓	euclidean_diff

Model	% Resolved	Avg. \$	Trajs	Org	Date	Agent
Claude 4.5 Opus (high reasoning)	76.80	\$0.75		AI	2026-02-17	2.0.0
Gemini 3 Flash (high reasoning)	75.80	\$0.36			2026-02-17	2.0.0
MiniMax M2.5 (high reasoning)	75.80	\$0.07			2026-02-17	2.0.0
Claude Opus 4.6	75.60	\$0.55		AI	2026-02-17	2.0.0
GPT-5-2 Codex	72.80	\$0.45			2026-02-19	2.0.0
GLM-5 (high reasoning)	72.80	\$0.53		Z	2026-02-17	2.0.0
GPT-5-2 (high reasoning)	72.80	\$0.47			2026-02-17	2.0.0
GPT 5.2 Codex	72.80	\$0.45			2026-02-19	2.0.0
Claude 4.5 Sonnet (high reasoning)	71.40	\$0.66		AI	2026-02-17	2.0.0
Kimi K2.5 (high reasoning)	70.80	\$0.15			2026-02-17	2.0.0
DeepSeek V3.2 (high reasoning)	70.00	\$0.45			2026-02-17	2.0.0
Gemini 3 Pro	69.60	\$0.96			2026-02-26	2.0.0
Claude 4.5 Haiku (high reasoning)	66.60	\$0.33		AI	2026-02-17	2.0.0
GPT-5 Mini	56.20	\$0.05			2026-02-17	2.0.0

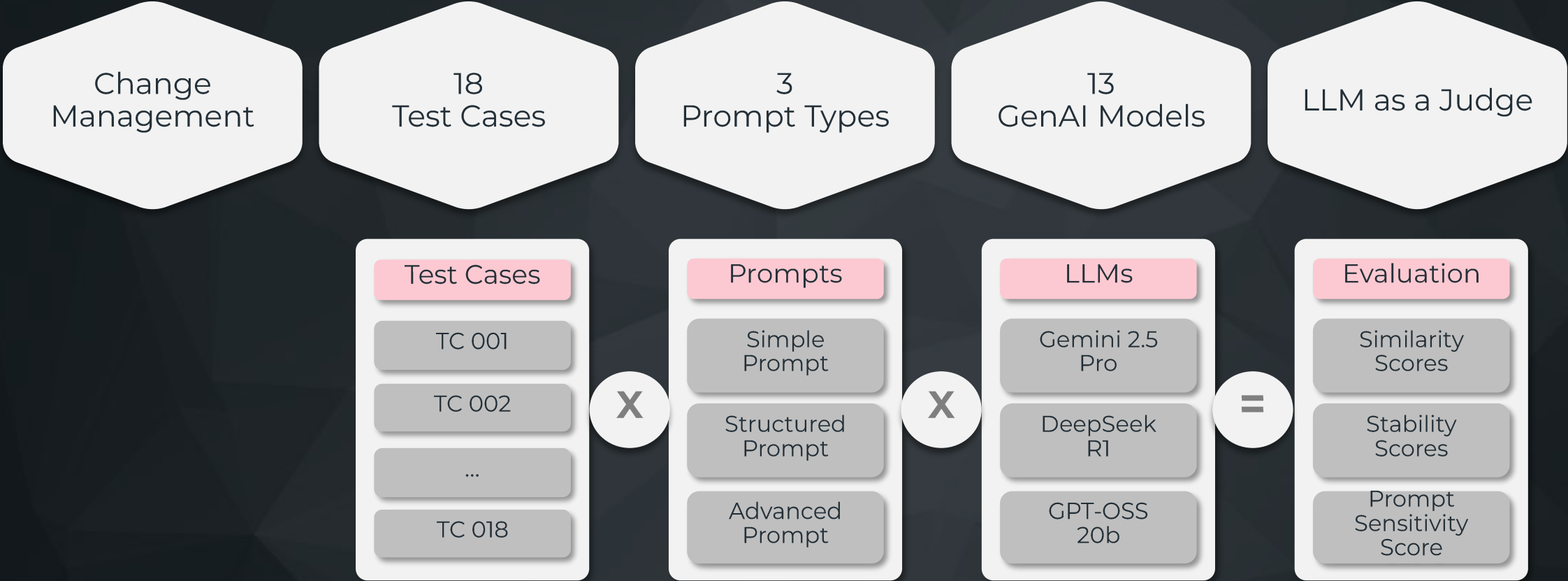
Benchmark for AI in Requirements Management ?



Benchmark for AI in Requirements Management ?



Our Solution : RM Bench



Test Case Example



Initial Input

The drone shall achieve a mean time between critical failures (MTBCF) of **at least 200 flight-hours**.



Change Logs

UAVs shall demonstrate a mean time between critical failures (MTBCF) of **no less than 300 flight-hours**, reflecting improved hardware reliability and customer expectations for longer uptime.



Prompt

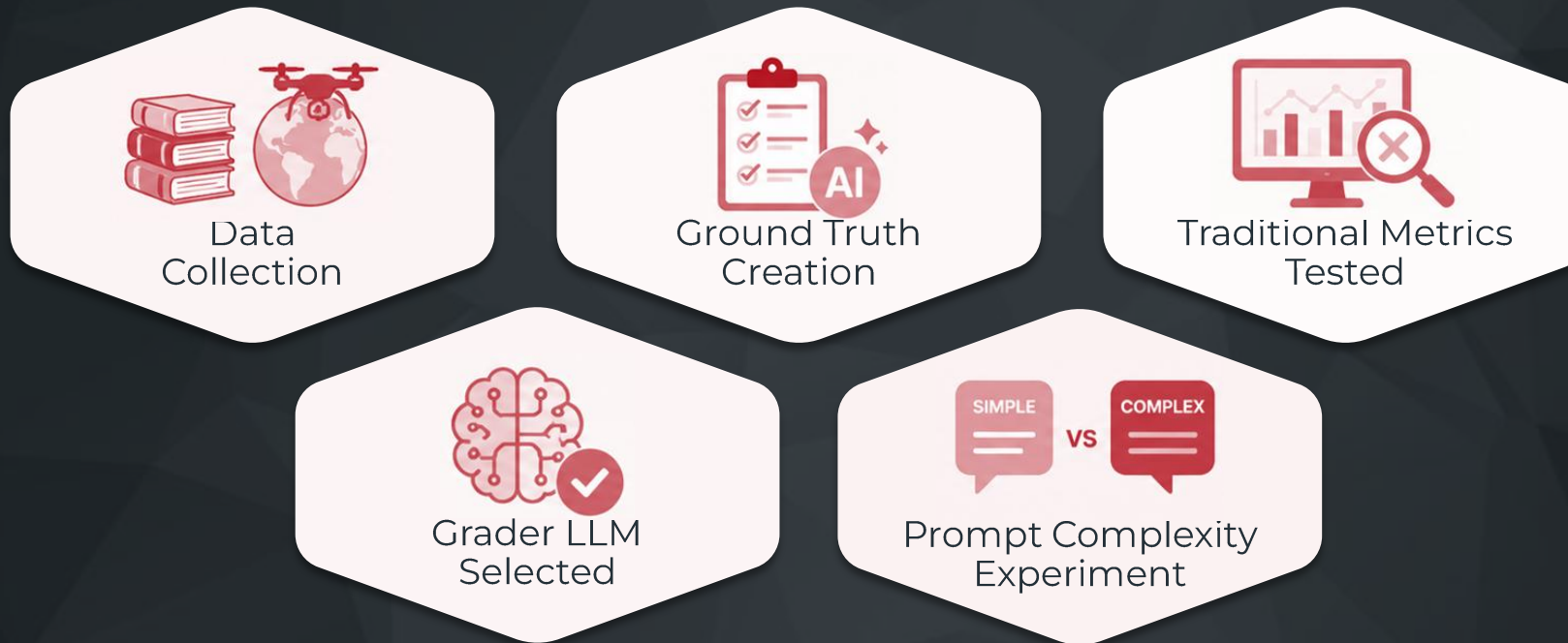
Update the list of requirements by changing, adding, or removing them exactly as told, and then give back the full updated list in the same format.



Golden Output

The drone shall achieve a mean time between critical failures (MTBCF) of **at least 300 flight-hours**.

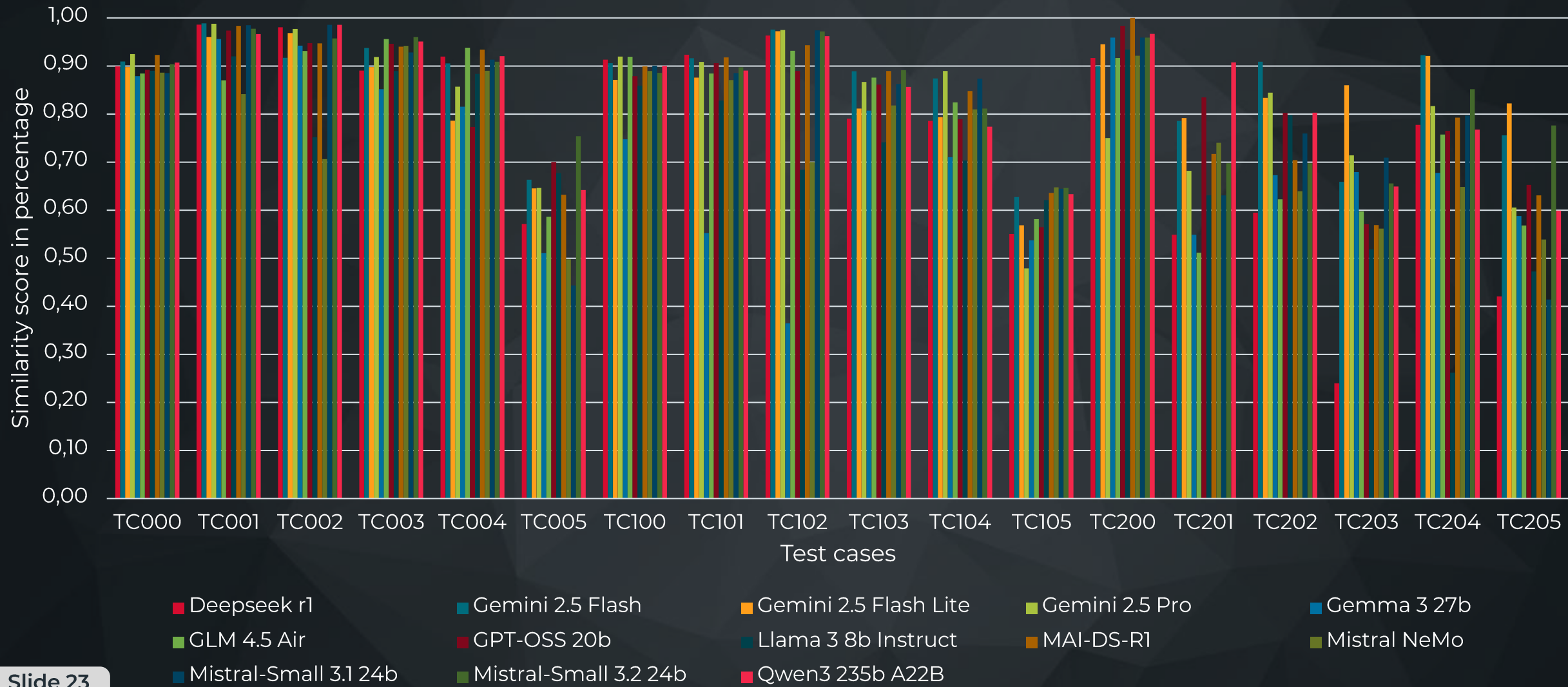
Experiment Setup



How it works ?

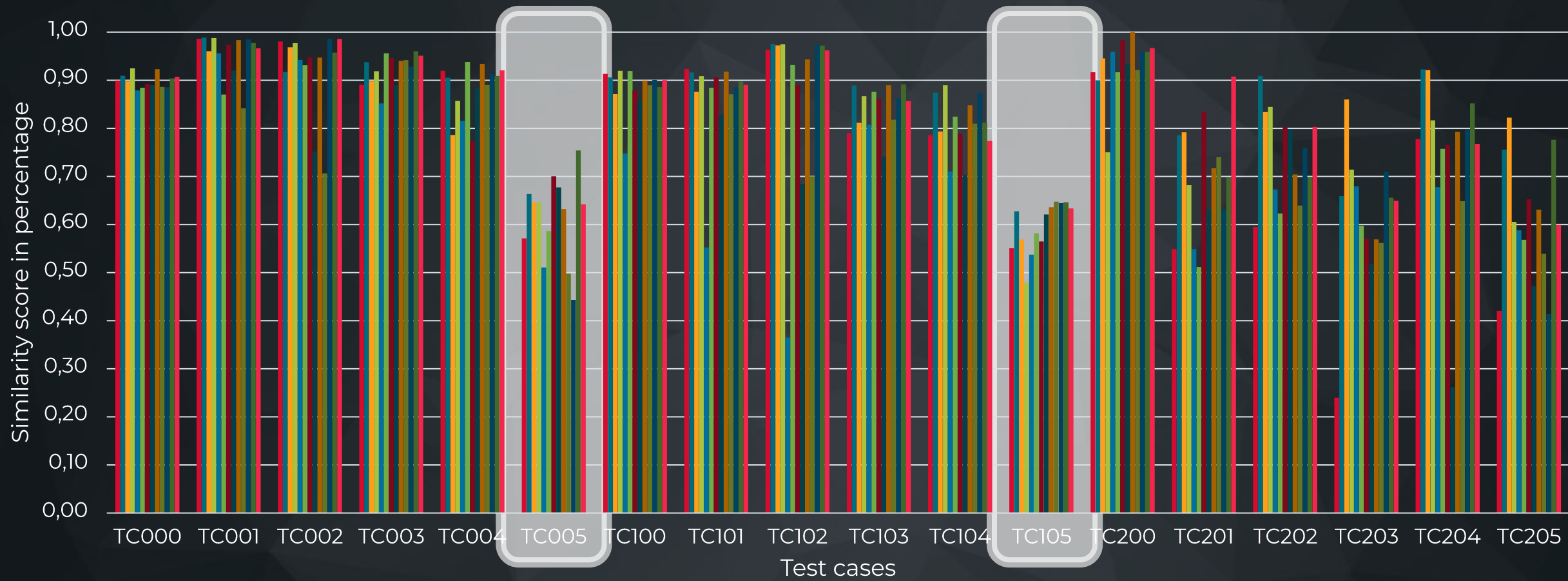


Our Findings





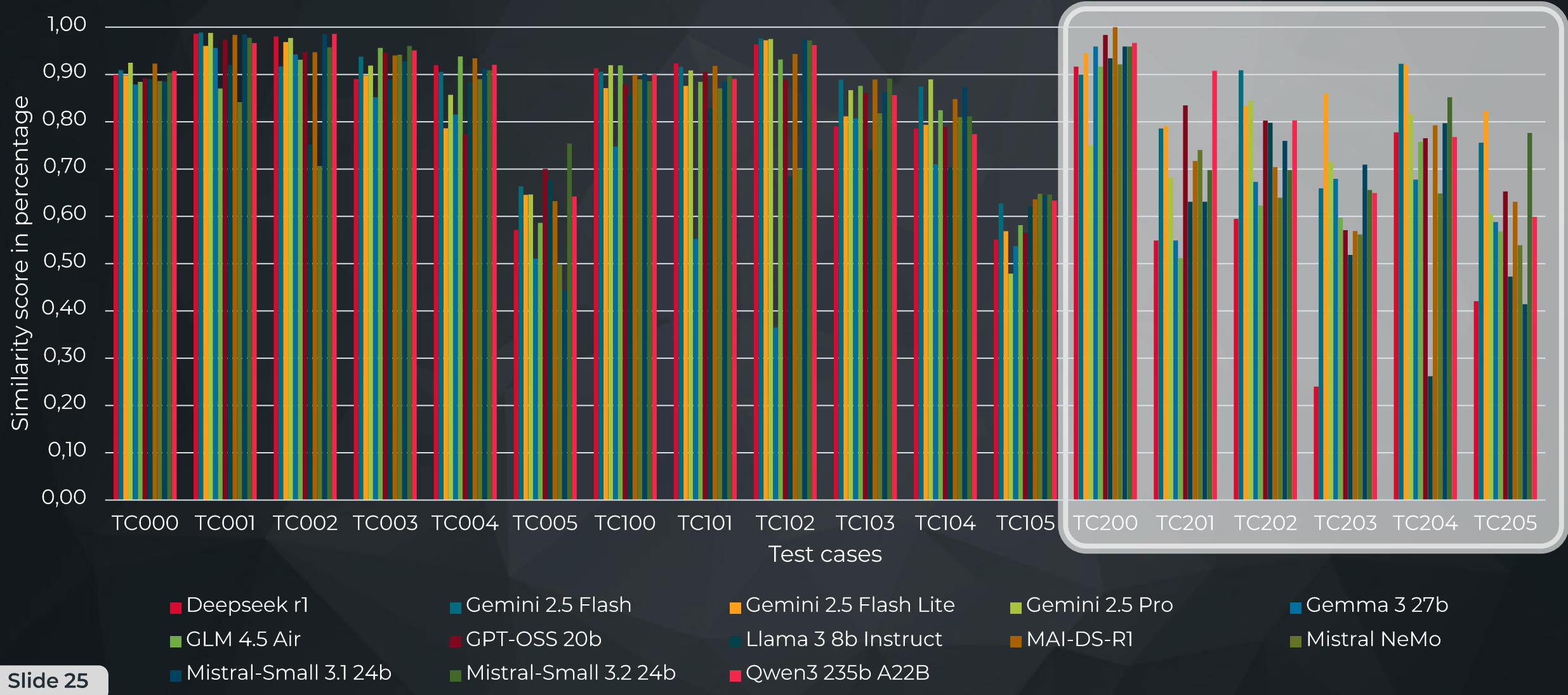
RM Bench : Our Findings



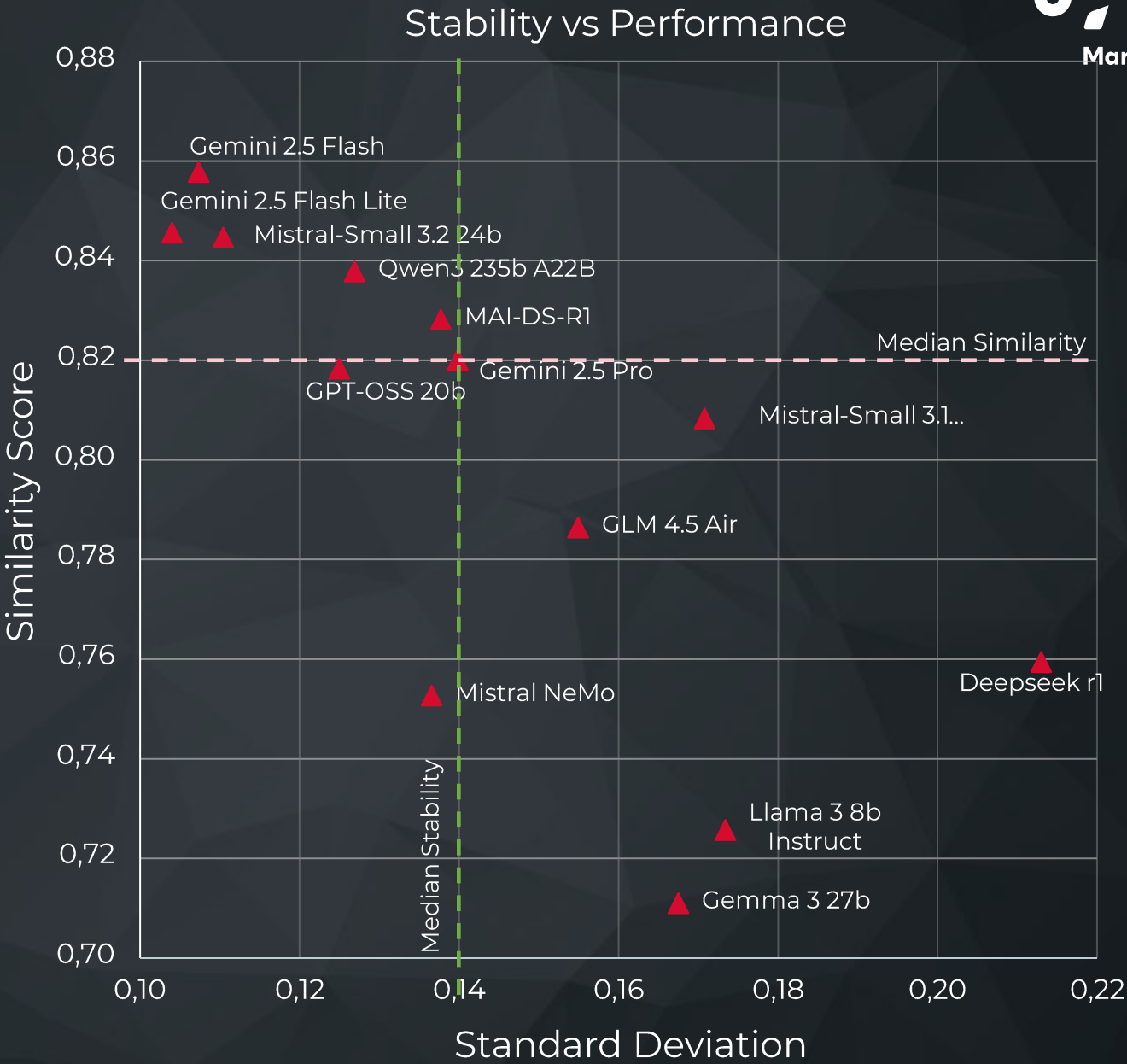
- Deepseek r1
- Gemini 2.5 Flash
- Gemini 2.5 Flash Lite
- Gemini 2.5 Pro
- Gemma 3 27b
- GLM 4.5 Air
- GPT-OSS 20b
- Llama 3 8b Instruct
- MAI-DS-R1
- Mistral NeMo
- Mistral-Small 3.1 24b
- Mistral-Small 3.2 24b
- Qwen3 235b A22B



RM Bench : Our Findings



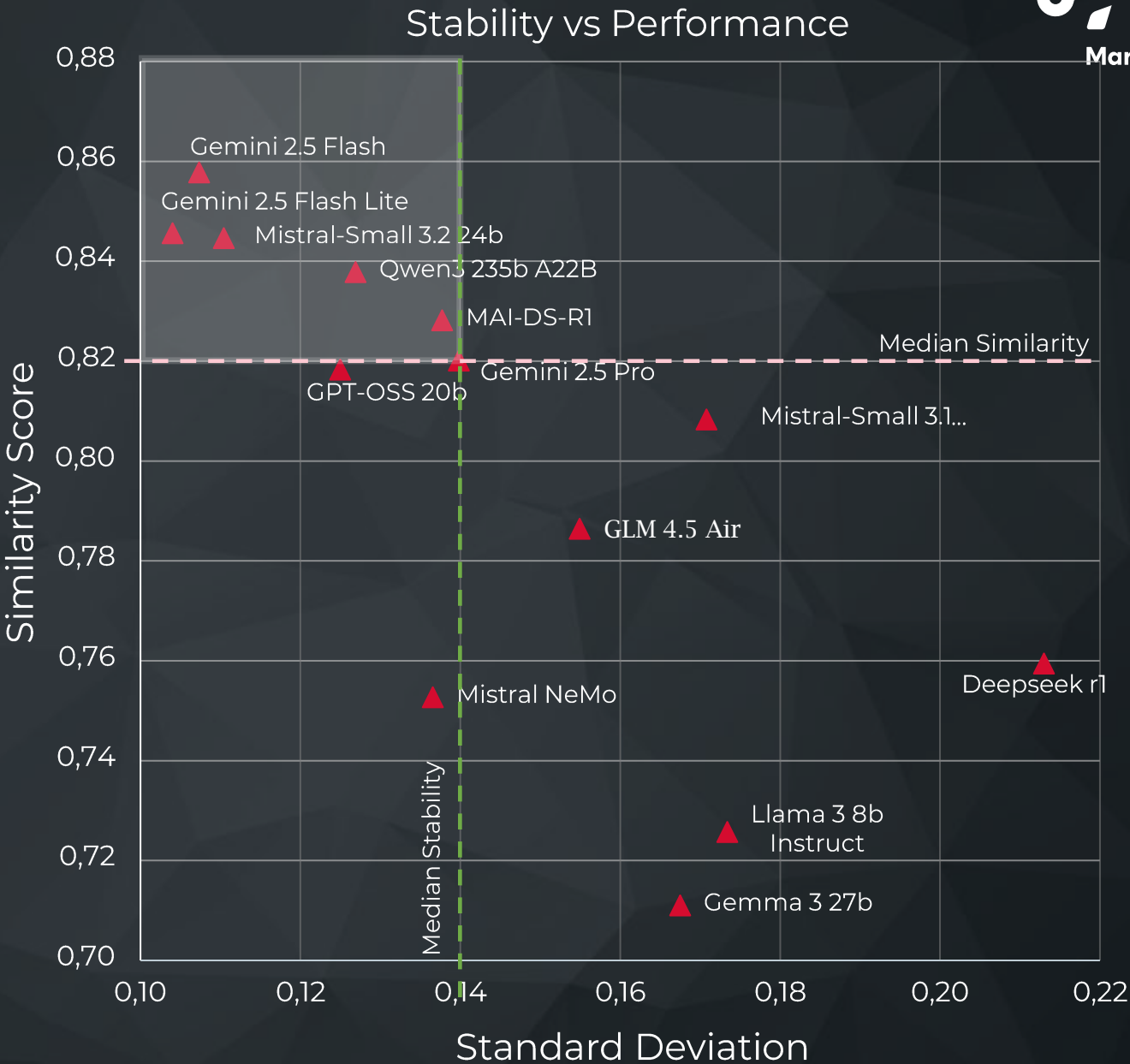
Overall Model Accuracy & Stability



Our Findings

Overall Model Accuracy & Stability

Top-performing models (e.g., Gemini 2.5 Flash, $\bar{s} \approx 0.86$) show high accuracy and low standard deviation (high stability).

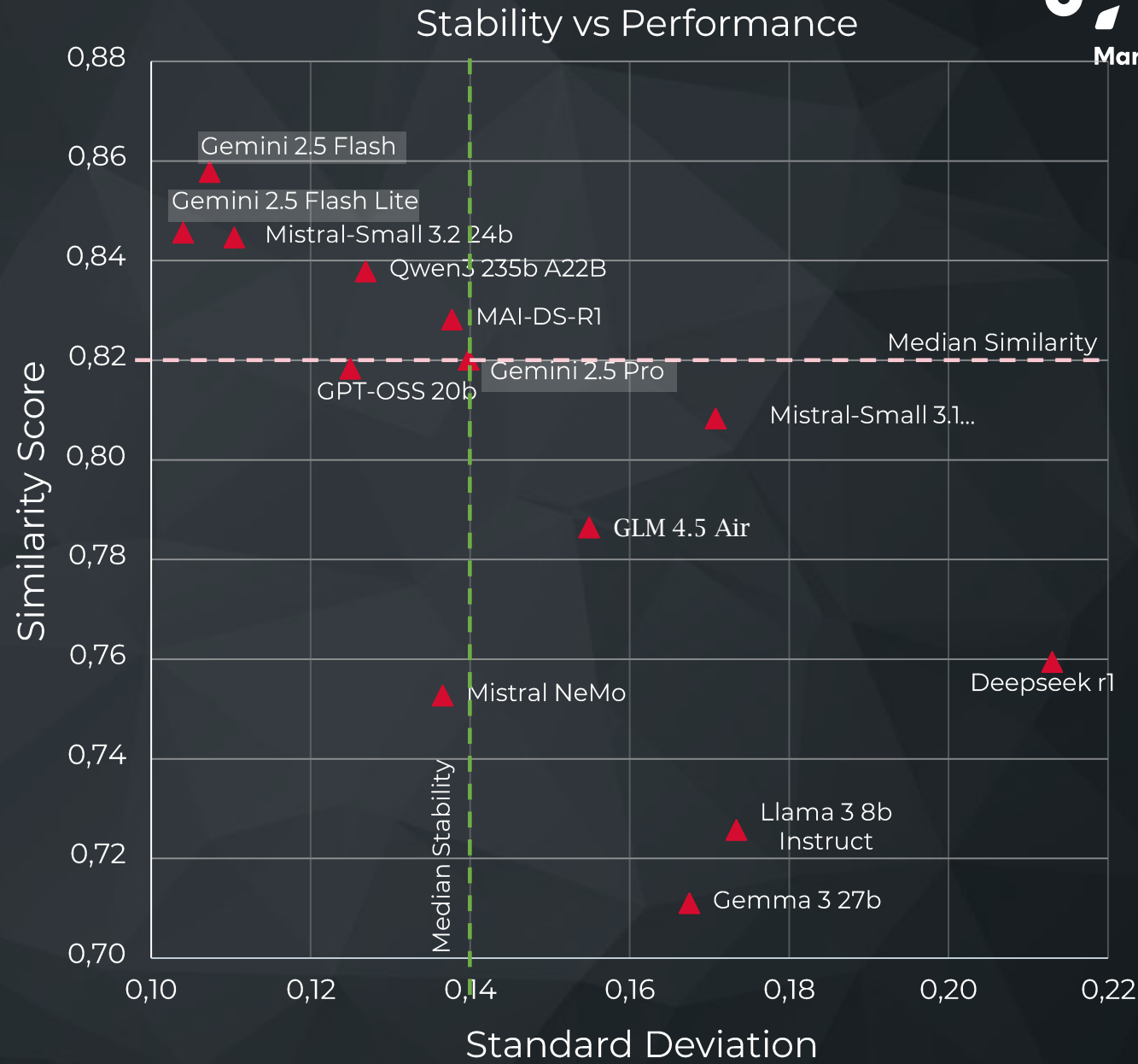


Our Findings

Overall Model Accuracy & Stability

Top-performing models (e.g., Gemini 2.5 Flash, $\bar{s} \approx 0.86$) show high accuracy and low standard deviation (high stability).

Gemini 2.5 Pro even being a better model than Gemini 2.5 Flash & Flash-Lite might be due to over-thinking.

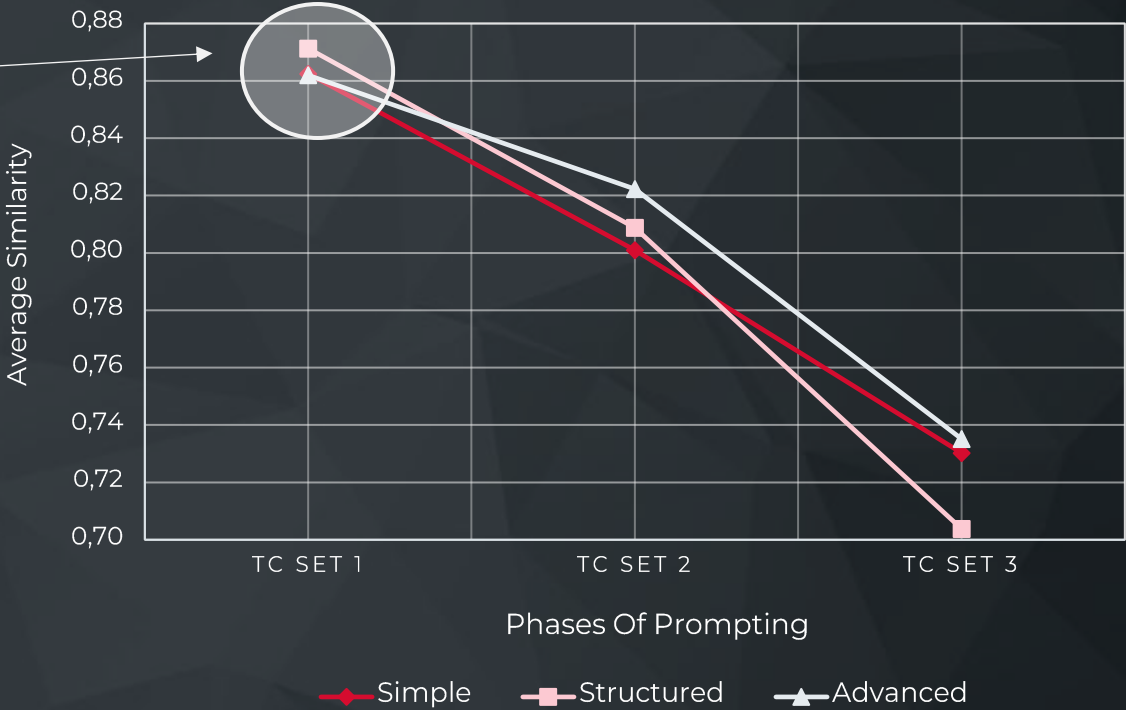


Performance Across Test Case Sets



TC Set 1 (Structured Tasks) achieved highest performance ($\bar{s} \approx 0.87$).

Prompt Type Trend Per TC Set



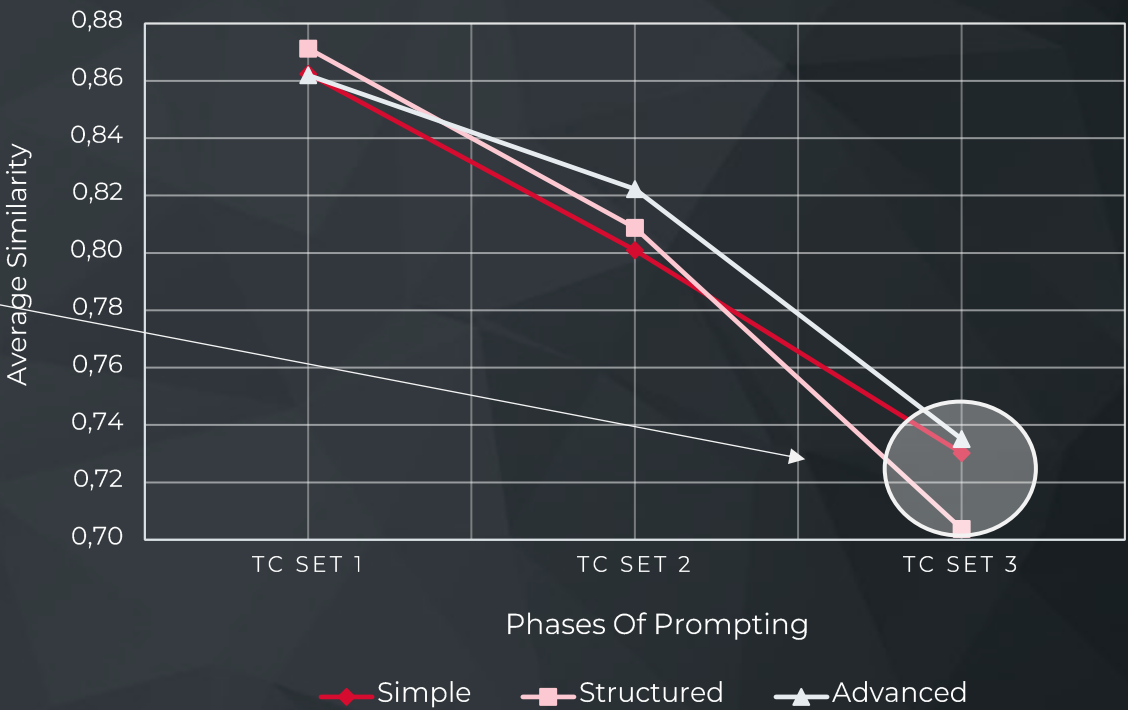
Performance Across Test Case Sets



TC Set 1 (Structured Tasks) achieved highest performance ($\bar{s} \approx 0.87$).

Performance dropped significantly in TC Set 3 (Semantic/Compliance Tasks) ($\bar{s} \approx 0.72$), confirming models struggle with **implicit regulatory semantics** and **contextual interpretation**.

Prompt Type Trend Per TC Set



Performance Across Test Case Sets



- TC Set 1 (Structured Tasks) achieved highest performance ($\bar{s} \approx 0.87$).
- Performance dropped significantly in TC Set 3 (Semantic/Compliance Tasks) ($\bar{s} \approx 0.72$), confirming models struggle with **implicit regulatory semantics** and **contextual interpretation**.
- Increasing complexity of tasks perform better with complex prompts.

Prompt Type Trend Per TC Set





Qualitative
Foundation



Competent But
Fragile



Highly Prompt
Sensitive

RM Bench

Limitations



GraderLLM
Bias



Input
Fragmentation

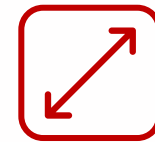


Limited
Scope

Future Work



Fine-tune RM
Based LLMs



Extend RM
Bench



Marketplace



Kontakt

nikhil.jha@ki-marktplatz.com
ruslan.bernijazov@ki-marktplatz.com