

Insper

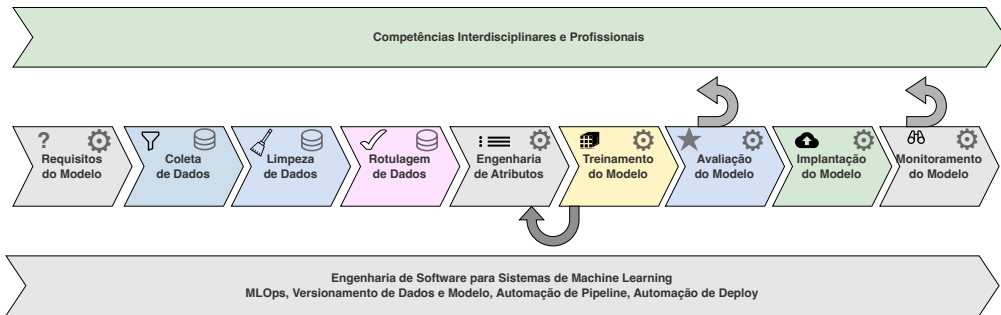
Agentes e o futuro da ciência de dados

da análise à orquestração, o novo papel do cientista de dados

INSPER

Fabício Barth, PhD

Projetos de ML tem uma forma de funcionamento bem definida



Agentes de IA estão transformando a execução de projetos de ML

- ▶ Agentes de IA aceleram a execução ao **gerar código**, consultas e estruturas para tarefas comuns, **reduzindo o tempo gasto** em implementações repetitivas.

Agentes de IA estão transformando a execução de projetos de ML

- ▶ Agentes de IA aceleram a execução ao **gerar código**, consultas e estruturas para tarefas comuns, **reduzindo o tempo gasto** em implementações repetitivas.
- ▶ Mudam a forma como cientistas de dados trabalham: **modo aceleração** (quando já sabem o próximo passo) e **modo exploração** (quando estão incertos sobre como proceder).

Agentes de IA estão transformando a execução de projetos de ML

- ▶ Agentes de IA aceleram a execução ao **gerar código**, consultas e estruturas para tarefas comuns, **reduzindo o tempo gasto** em implementações repetitivas.
- ▶ Mudam a forma como cientistas de dados trabalham: **modo aceleração** (quando já sabem o próximo passo) e **modo exploração** (quando estão incertos sobre como proceder).
- ▶ Reduzem a carga cognitiva em tarefas repetitivas, **liberando tempo para análises de maior valor**, como interpretação de resultados, formulação de hipóteses e validação crítica de modelos.

Agentes de IA estão transformando a execução de projetos de ML

- ▶ Agentes de IA aceleram a execução ao **gerar código**, consultas e estruturas para tarefas comuns, **reduzindo o tempo gasto** em implementações repetitivas.
- ▶ Mudam a forma como cientistas de dados trabalham: **modo aceleração** (quando já sabem o próximo passo) e **modo exploração** (quando estão incertos sobre como proceder).
- ▶ Reduzem a carga cognitiva em tarefas repetitivas, **liberando tempo para análises de maior valor**, como interpretação de resultados, formulação de hipóteses e validação crítica de modelos.
- ▶ Deslocam o cientista de dados de escrever tudo manualmente para **prompting**, **revisão** e **orquestração** das etapas do projeto.

Mas ainda existem desafios no uso de agentes de IA em projetos de ML

- ▶ **Alucinações e erros:** agentes geram código inseguro, dependências inexistentes e replicam práticas vulneráveis.

Mas ainda existem desafios no uso de agentes de IA em projetos de ML

- ▶ **Alucinações e erros:** agentes geram código inseguro, dependências inexistentes e replicam práticas vulneráveis.
- ▶ **Necessidade de mais Revisão Humana:** os ganhos de velocidade aparentes podem ser anulados por verificações e correções posteriores.

Mas ainda existem desafios no uso de agentes de IA em projetos de ML

- ▶ **Alucinações e erros:** agentes geram código inseguro, dependências inexistentes e replicam práticas vulneráveis.
- ▶ **Necessidade de mais Revisão Humana:** os ganhos de velocidade aparentes podem ser anulados por verificações e correções posteriores.
- ▶ **Super Dependência e perda de Habilidades:** usuários que aceitam sugestões sem compreendê-las ficam sujeitos a **viés de automação**.

Mas ainda existem desafios no uso de agentes de IA em projetos de ML

- ▶ **Alucinações e erros:** agentes geram código inseguro, dependências inexistentes e replicam práticas vulneráveis.
- ▶ **Necessidade de mais Revisão Humana:** os ganhos de velocidade aparentes podem ser anulados por verificações e correções posteriores.
- ▶ **Super Dependência e perda de Habilidades:** usuários que aceitam sugestões sem compreendê-las ficam sujeitos a **viés de automação**.
- ▶ **Raciocínio analítico complexo:** agentes são consistentemente fortes em tarefas estruturadas e repetitivas, mas fracos em raciocínio complexo e operação autônoma sem supervisão.

Mas ainda existem desafios no uso de agentes de IA em projetos de ML

- ▶ **Alucinações e erros:** agentes geram código inseguro, dependências inexistentes e replicam práticas vulneráveis.
- ▶ **Necessidade de mais Revisão Humana:** os ganhos de velocidade aparentes podem ser anulados por verificações e correções posteriores.
- ▶ **Super Dependência e perda de Habilidades:** usuários que aceitam sugestões sem compreendê-las ficam sujeitos a **viés de automação**.
- ▶ **Raciocínio analítico complexo:** agentes são consistentemente fortes em tarefas estruturadas e repetitivas, mas fracos em raciocínio complexo e operação autônoma sem supervisão.
- ▶ **Privacidade e governança:** riscos estruturais que persistem independentemente de como o agente é usado. Exigem **novas políticas**, controles e supervisão organizacional.

Referências

- Peng et al. (2023) The Impact of AI on Developer Productivity: Evidence from GitHub Copilot
- Ullsnes et al. (2024) Transforming Software Development with Generative AI: Empirical Insights on Collaboration and Workflow
- Barke et al. (2022) Grounded Copilot: How Programmers Interact with Code-Generating Models
- McNutt et al. (2023) On the Design of AI-powered Code Assistants for Notebooks
- Zhou & Li (2023) A Case Study on Scaffolding Exploratory Data Analysis for AI Pair Programmers
- Gu et al. (2023) How Do Data Analysts Respond to AI Assistance? A Wizard-of-Oz Study
- Gupta (2025) The Rise of AI Copilots: Redefining Human–Machine Collaboration in Knowledge Work
- Bird et al. (2022) Taking Flight with Copilot
- Kudriavtseva et al. (2025) My Code Is Less Secure with Gen AI: Surveying Developers' Perceptions of the Impact of Code Generation Tools on Security