

Capítulo

2

Ferramentas de Inteligência Artificial Para Suporte ao Processo de Pesquisa em Sistemas de Informação

Artificial Intelligence Tools to Support the Research Process in Information Systems

Fábio Manoel França Lobato, Matheus Lima do Couto, Antonio Fernando Lavareda Jacob Jr., René Vieira Santin, Solange Oliveira Rezende, Ricardo Marcondes Marcacini

Abstract

The rise of Artificial Intelligence (AI), driven by Large Language Models (LLMs), is transforming scientific research, generating an exponential volume of publications and a new ecosystem of support tools. However, many researchers remain unaware of how to integrate these technologies effectively and ethically into their workflows. This introductory mini-course aims to fill this gap, equipping the Information Systems (IS) community with tools and best practices to support the entire AI research lifecycle. With an eminently practical, hands-on approach, participants will explore topics ranging from fundamental concepts of LLMs to the use of platforms such as Research Rabbit for literature retrieval and NotebookLM for analysis. Prompt engineering techniques for using LLMs as assistants for ideation, scientific writing, and reproducibility will be covered. A central focus of the course will be the discussion of ethical and transparency challenges, culminating in the collaborative construction of a guide to best practices. The aim is to enable participants to apply AI productively, critically, and with integrity in their research, supported by accessible materials that reflect the field's inter and transdisciplinary nature.

Resumo

A ascensão da Inteligência Artificial (IA), impulsionada por Modelos de Linguagem de Grande Escala (LLMs), está transformando a pesquisa científica, gerando um volume exponencial de publicações e um novo ecossistema de ferramentas de apoio. No entanto,

muitos pesquisadores ainda desconhecem como integrar essas tecnologias de forma eficaz e ética ao seu fluxo de trabalho. Este minicurso introdutório visa preencher essa lacuna, instrumentalizando a comunidade de Sistemas de Informação (SI) com ferramentas e boas práticas para apoiar todo o ciclo de vida da pesquisa com IA. Com uma abordagem eminentemente prática e hands-on, os participantes explorarão desde conceitos fundamentais de LLMs até o uso de plataformas como Research Rabbit para recuperação de literatura e NotebookLM para análise. Serão abordadas técnicas de engenharia de prompts para utilizar LLMs como assistentes na ideação, na escrita científica e na reprodutibilidade. Um eixo central do minicurso será a discussão dos desafios éticos e de transparência, culminando na construção colaborativa de um guia de boas práticas. Espera-se capacitar os participantes a aplicar a IA de forma produtiva, crítica e íntegra em suas investigações, com material de apoio acessível que contemple o caráter inter e transdisciplinar da área de SI.

2.1. Introdução

A Inteligência Artificial, impulsionada pelos avanços em Modelos de Linguagem de Grande Escala, também conhecidos pela terminologia em inglês *Large Language Models* (LLMs), está revolucionando inúmeros setores. A pesquisa científica figura entre os campos mais profundamente transformados [Mishra et al. 2024]. Essas tecnologias inauguraram uma nova era na produção do conhecimento, oferecendo capacidades sem precedentes para sintetizar informações complexas, automatizar tarefas de revisão de literatura, gerar novas hipóteses [Si et al. 2024] e acelerar a análise de dados [Chen et al. 2025, Gottweis et al. 2025b].

Este novo paradigma já traz consigo um impacto direto e um desafio crescente: o aumento exponencial da produção técnico-científica. Tanto que a comunidade científica tem registrado números recordes de submissões, o que afeta diretamente a carga de trabalho nas revisões. A título de exemplo, o *Encontro Nacional de Inteligência Artificial e Computacional* (ENIAC) de 2025 solicitou pelo menos 10 revisões a cada membro do seu comitê de programa. O Workshop sobre Educação em Computação (WEI) de 2025 experimentou um aumento de 197% nos últimos dois anos, tendo recebido 347 submissões em 2025 contra 176 em 2023.

Nesse contexto, floresceu um ecossistema de ferramentas de IA projetadas especificamente para o fazer científico, uma profusão de modelos de IA e sistemas inteligentes cobrindo tarefas como a compreensão de textos científicos, automação da revisão da literatura, construção e aperfeiçoamento de desenhos de pesquisa, recomendação de locais de publicação, suporte e semiautomatização da escrita acadêmica, até mesmo a revisão por “pares”, apenas para mencionar alguns [Chen et al. 2025]. No campo de Sistemas de Informação (SI), destacam-se, em particular, aplicações de IA voltadas ao apoio ao enquadramento da pesquisa em *frameworks* teóricos e epistemológicos, ao alinhamento do desenho de pesquisa aos grandes desafios contemporâneos e ao fortalecimento da reprodutibilidade experimental, as quais apresentam elevado potencial de exploração. Neste sentido, é urgente que a comunidade seja instrumentalizada com uma capacitação que:

- Forneça uma visão geral das ferramentas de IA para suporte à pesquisa científica atualmente existentes;

- Apresente aplicações de modelos de IA generativa que demonstrem seu potencial para auxiliar no fazer científico;
- Promova uma discussão aprofundada sobre as implicações éticas, de transparência e de reprodutibilidade associadas ao uso da IA na pesquisa;
- Ofereça demonstrações práticas e casos de uso de como essas ferramentas podem ser aplicadas para resolver desafios específicos da pesquisa em SI, fortalecendo o rigor metodológico da comunidade;
- Prepare o novo pesquisador para o futuro presente da Pesquisa em SI, promovendo a fluência no uso de ferramentas de IA, competência necessária neste novo mundo.

Conforme apresentado, apesar do imenso potencial para aumentar a produtividade, o rigor e o alcance das investigações, boa parte da comunidade científica ainda desconhece este ferramental ou não se sente segura para integrá-lo de forma eficaz e ética ao seu fluxo de trabalho. Visando preencher essa lacuna, o presente minicurso propõe uma jornada introdutória e prática. O intuito é capacitar os participantes não apenas a conhecer, mas também a aplicar criticamente as principais tecnologias de IA disponíveis, discutindo boas práticas e os princípios éticos indispensáveis à plena e íntegra incorporação da IA às pesquisas em SI. Neste sentido, este minicurso tem como objetivo instrumentalizar pessoas pesquisadoras com os principais avanços da Inteligência Artificial como suporte ao processo de pesquisa científica em SI, apresentando ferramentas, boas práticas e princípios aplicáveis ao seu cotidiano de pesquisa. Ao final deste minicurso, espera-se que o participante seja capaz de:

- Compreender os conceitos fundamentais sobre Modelos de Língua de Grande Escala e seu impacto no ciclo de vida da pesquisa científica;
- Aplicar ferramentas de IA para otimizar tarefas de busca, recuperação, sumarização e análise de literatura científica;
- Desenvolver *prompts* eficazes para utilizar LLMs como assistentes na ideação, revisão de textos, análise de dados e construção de artefatos;
- Avaliar criticamente as potencialidades, limitações e riscos associados a cada ferramenta;
- Identificar e debater os principais desafios éticos associados ao uso de ferramentas de IA no contexto da pesquisa em Sistemas de Informação.

O escopo do curso é de caráter **introdutório**, até pela sua duração prevista de quatro horas, se possível, dividida em dois momentos de duas horas cada. Os **requisitos técnicos** são básicos; sugere-se que cada participante tenha um notebook com conexão à internet. O único software necessário será um navegador web atualizado (*e.g.*, Chrome, Firefox etc.) e contas gratuitas nas ferramentas, entre outras a serem indicadas. O conteúdo programático prevê seis eixos principais, a saber:

1. Uma visão geral de modelos de língua de grande escala e conceitos básicos de cientometria;

2. Ferramentas para a recuperação de artigos (*e.g.*, Research Rabbit);
3. Ferramenta para sumarização e análise de artigos, incluindo a criação de podcasts (*e.g.*, NotebookLM);
4. Explorando LLMs como um co-cientista (*e.g.*, Gemini, GPT e DeepSeek);
5. Engenharia de *prompt* para a revisão, reprodutibilidade e construção de artefatos;
6. Aspectos éticos e de responsabilidade no uso de ferramentas baseadas em IA para a pesquisa em SI.

O primeiro é o único de caráter mais teórico, enquanto os quatro seguintes são de caráter mais prático. O último visa a engajar os participantes na construção conjunta de um *guide-line* para o uso ético e responsável de ferramentas de IA. Nesse sentido, este capítulo tem como objetivo possibilitar a replicação do minicurso, ampliando seu alcance e servindo de subsídio à incorporação de conceitos, ferramentas e práticas em cursos similares, tais como disciplinas de metodologia científica, ética em pesquisa, produção acadêmica, entre outras. Todos os *prompts* e materiais suplementares estão publicamente disponíveis no GitHub: https://github.com/fabiolobato/cursoiapesquisa_sbsi, sob licença CC BY-NC 4.0.

O capítulo encontra-se organizado da seguinte maneira: na introdução foram apresentadas a visão geral do curso, a motivação, as habilidades e competências a serem desenvolvidas, bem como o conteúdo programático. A Seção 2.2 apresenta a fundamentação teórica, abordando o histórico dos Modelos de Linguagem de Grande Escala e seu panorama atual, fornecendo subsídios para a compreensão dos aspectos subjacentes ao ferramental utilizado. A Seção 2.3 apresenta um panorama do uso da IA como co-cientista. O panorama das ferramentas selecionadas é apresentado na Seção 2.4. Os exemplos de uso são apresentados e discutidos na Seção 2.5, com detalhamento dos *prompts* e de exemplos de saída. Aspectos didáticos e éticos são discutidos nas Seções 2.6 e 2.7. Por fim, as considerações finais são apresentadas na Seção 2.8.

2.2. Modelos de Linguagem de Grande Escala

Apesar de terem se popularizado recentemente, os LLMs são o resultado de um desenvolvimento longo e contínuo na modelagem de linguagem e na PLN. Já na década de 1950, trabalhos inspirados na teoria da informação analisavam modelos simples baseados em sequências de palavras, o que ajudou a consolidar a modelagem estatística da linguagem como base para tarefas como reconhecimento de fala, tradução automática e recuperação de informação [Toosi et al. 2021, Liddy 2001].

Formalmente, um **modelo de língua** (ou **modelo de linguagem**) define uma distribuição de probabilidade sobre sequências de *tokens*. Seja uma sequência $x_{1:T} = (x_1, x_2, \dots, x_T)$, um modelo de linguagem estima o próximo *token* com base na Equação 1:

$$p(x_{1:T}) = \prod_{t=1}^T p(x_t | x_{<t}), \quad (1)$$

em que $x_{<t}$ representa o histórico anterior ao *token* x_t . Em modelos *causais* (como os GPT), a tarefa central de treino é prever o próximo *token* a partir do histórico. A seguir, um exemplo simples que ilustra a ideia. Se durante o treino o modelo observa sentenças do tipo: “A tradução de mesa bonita para inglês é *beautiful table*”. Então, ao receber o prefixo “A tradução de mesa bonita para inglês é”, ele tende a atribuir alta probabilidade à continuação “*beautiful table*”. Esse comportamento decorre diretamente da fatoração acima: a cada passo, o modelo escolhe o próximo *token* com base no contexto em que o *token* aparece.

A base conceitual das redes neurais artificiais usadas hoje em LLMs remonta ao trabalho de McCulloch e Pitts, de 1943, que propôs um modelo lógico de neurônio. Isso ocorreu antes mesmo do termo *Inteligência Artificial* ser cunhado na Conferência de Dartmouth, em 1956. Ainda nesse período, ideias centrais, como o aprendizado por reforço, surgem com Arthur Samuel (também em 1956), antecipando técnicas que mais tarde seriam reutilizadas no alinhamento de modelos (por exemplo, RLHF) [Toosi et al. 2021]. Em paralelo, o Processamento de Linguagem Natural (PLN) começa a se estruturar como área, com destaque para a tradução automática como tarefa motivadora [Liddy 2001, Jones 1994]. Nessa época, também surgem sistemas iniciais de recuperação de informação baseados em *thesaurus*, que já apontavam para a centralidade da representação de texto e da busca [Liddy 2001]. A seguir, são brevemente apresentados quatro estágios de desenvolvimento dos LLMs, que exploram os principais até o momento.

2.2.1. Estágio 1: modelos estatísticos de linguagem e limites de contexto

Os primeiros modelos eram, em geral, diretos e frágeis: mapeavam palavras e aplicavam poucas regras de reorganização, ignorando o contexto amplo [Joseph et al. 2016]. Esse tipo de limitação aparece claramente em modelos estatísticos clássicos, como os *n-gramas*. Neles, assume-se uma dependência curta entre *tokens*, o que simplifica o cálculo, mas reduz a capacidade de capturar relações que dependem de contextos mais amplos. Além disso, como as probabilidades são estimadas por contagens em que *tokens* ocorrem juntos em alguma janela de contexto, ocorre o efeito da **esparsidade**, pelo qual, na prática, muitos *n-gramas* possíveis não aparecem no treino.

Na mesma linha histórica, a influência dos estudos de Chomsky (1957) marcou a busca por estruturas formais da linguagem, enquanto métodos estatísticos ganharam força no reconhecimento de fala, criando, ao longo dos anos, uma tensão entre abordagens simbólicas e probabilísticas [Liddy 2001].

2.2.2. Estágio 2: modelos neurais e *embeddings*

Com o aumento de dados e de poder computacional, modelos neurais passam a substituir contagens por representações contínuas. Neste tipo de modelo, o objetivo é aprender **embeddings** de palavras, representações vetoriais que permitem computar a proximidade entre palavras, e usar diferentes arquiteturas de redes neurais para combinar o contexto e prever o próximo *token*. Assim, é possível reduzir o problema de esparsidade que existia nos *n-gramas*, pois mesmo sem ver uma sequência específica, o modelo pode generalizar por similaridade no espaço vetorial.

Em tarefas sequenciais, um marco foi o uso de *Recurrent Neural Networks* (RNNs)

a partir da década de 1970 e, mais tarde, de arquiteturas mais estáveis, como o *Long Short-Term Memory* (LSTM), que controlam o fluxo de informação (o que guardar e o que descartar) para lidar melhor com dependências longas [Mienye et al. 2024]. Em 2014, a atenção aplicada a RNNs mudou esse cenário ao permitir que o decodificador *revisite* toda a sentença de entrada em cada etapa de geração, com enfoque automático nas partes mais relevantes [Bahdanau et al. 2014].

2.2.3. Estágio 3: modelos pré-treinados e ajuste fino

Os **modelos pré-treinados** consolidaram uma mudança de estratégia na qual, em vez de treinar um modelo do zero para cada tarefa, realiza-se um **pré-treino** em larga escala com dados não rotulados (por objetivos de auto-supervisão, como a previsão de *tokens*) e, depois, um **ajuste** (*fine-tuning*) para tarefas específicas com poucos dados rotulados. É importante notar que modelos pré-treinados podem ser baseados em RNNs ou em *Transformers*, a última sendo a arquitetura mais utilizada atualmente.

Um marco inicial foi o *Embeddings from Language Models* (**ELMo**), que popularizou **embeddings contextuais** [Peters et al. 2018]. Nele, a representação de uma palavra passa a depender da frase inteira, e não apenas de um vetor fixo por vocabulário. Na prática, isso significa que a mesma palavra pode ter *embeddings* diferentes conforme o contexto (por exemplo, “banco” em finanças versus “banco” como assento). O ELMo foi treinado como um modelo de linguagem (com redes recorrentes) e suas camadas internas passaram a servir como fonte geral de *embeddings*.

Em seguida, o **BERT** [Devlin et al. 2019] consolidou o estágio de modelos pré-treinados ao explorar *Transformers* e ao introduzir um objetivo de pré-treino baseado em *Masked Language Modeling* (MLM). Neste caso, o modelo aprende a reconstruir *tokens* previamente mascarados (propositalmente) a partir do contexto. Após o pré-treino, o BERT pode ser ajustado para outras tarefas, muitas vezes com ganhos significativos mesmo com poucos exemplos rotulados.

2.2.4. Estágio 4: LLMs como modelos fundacionais

LLMs também são **pré-treinados**, mas o estágio de LLM se diferencia dos modelos anteriores por um conjunto de fatores. Primeiro, na **(i) escala e regime de treino**, pois LLMs tipicamente operam com ordens de grandeza maiores em parâmetros, dados e computação. Essa escala se tornou viável com *Transformers*, que removem a dependência sequencial no treino e permitem paralelismo massivo [Vaswani et al. 2017]. Em segundo lugar, nas **(ii) capacidades emergentes no uso de prompts**. Em LLMs, surgiu mais fortemente o conceito de *in-context learning*, em que o modelo aprende uma tarefa a partir de exemplos no próprio *prompt*, sem atualizar pesos. Isso muda o modo de uso em relação aos modelos anteriores, pois reduz a dependência de *fine-tuning* em cada tarefa e amplia a reutilização imediata.

As LLMs, como modelos funcionais, também permitem a etapa de **pós-treino** voltada a seguir instruções (*instruction tuning*) e ao alinhamento por *feedback* humano, o que melhora a utilidade, a consistência e a adequação na interação [Wu et al. 2023]. Esse passo é importante para transformar um bom preditor de *tokens* em um assistente que responde a tarefas de forma mais controlada. Também se observa que LLMs são

frequentemente usados como componentes centrais de sistemas maiores (por exemplo, assistentes e agentes), com decisões sobre quando recuperar informação, quando chamar outras ferramentas e como organizar os passos intermediários.

Historicamente, a família GPT ilustra bem essa transição. O GPT mostrou o potencial de um *Transformer decoder-only* como um modelo generativo geral [Radford et al. 2018]. O GPT-2 (2019) ampliou a escala e os dados e passou a operar como um modelo mais geral, embora ainda apresente instabilidades em longas gerações [Radford et al. 2019]. O GPT-3 (2020) consolidou *few-shot learning* e ampliou a capacidade de generalização com base em exemplos no *prompt* [Brown et al. 2020]. Em 2022, a popularização do ChatGPT reforçou o papel do pós-treino por instruções para uso interativo [Wu et al. 2023]. Em paralelo, modelos como PaLM e versões especializadas (por exemplo, voltadas à matemática e às ciências) reforçaram a tendência de combinar escala com especialização e técnicas de raciocínio guiado [Annepaka and Pakray 2025].

No cenário atual, a consolidação do uso de LLMs por meio de APIs e restrições de uso incentivaram o surgimento de iniciativas de código aberto e de modelos com licenças mais permissivas, ampliando o acesso a pesos, à reprodutibilidade e à adaptação em cenários acadêmicos e industriais [Annepaka and Pakray 2025]. Ao mesmo tempo, surgiram especializações por domínio (*e.g.*, finanças e saúde) e por modalidade (por exemplo, texto+imagem), o que ampliou o escopo de aplicações e as exigências de avaliação, governança e mitigação de vieses [Peykani et al. 2025, Wang et al. 2025].

Outro ponto comum atualmente é a presença de limitações importantes, como alucinações, sensibilidade ao *prompt* e risco de desatualização. Uma abordagem prática para reduzir esses problemas é a **geração aumentada por recuperação**, mais conhecida como *Retrieval-Augmented Generation* (RAG), em que o modelo recebe, junto ao *prompt*, evidências recuperadas de fontes confiáveis e atuais. Em termos de sistema, a RAG incorpora no *prompt* um conjunto explícito de documentos relevantes, obtidos por meio de uma busca ou de alguma ferramenta, o que favorece respostas mais ancoradas [Peykani et al. 2025]. Em cenários mais complexos, LLMs são integrados a **agentes**, que planejam passos, consultam ferramentas, buscam dados e iteram com o usuário, tornando o modelo parte de um ciclo de decisão e execução [OpenAI 2025, DeepMind 2025].

Por fim, destaca-se que muitas das aplicações atuais que se convencionou chamar de *IA como co-cientista*, a serem apresentadas a seguir, são instâncias de LLMs operando em diferentes configurações. Por exemplo, um “co-cientista” pode ser, na prática, desde um modo RAG, em que recupera trechos, evidências ou documentos completos a partir de bases científicas para melhorar seu *prompt* e seu contexto, até um modo *agente*, em que busca ativamente informações em fontes heterogêneas usando ferramentas (como *browsers* e outras APIs), combina resultados e interage com o usuário ao longo de vários passos. Também será apresentado que, além dessas arquiteturas recentes, diversas capacidades que já eram típicas continuam a ser usadas nesse cenário, como a extração e a organização de informações e a sumarização multidocumento (*e.g.*, artigos científicos).

2.3. IA como o co-cientista

O uso de modelos de linguagem e de ferramentas de IA aplicados ao contexto da pesquisa acadêmica já se encontra amplamente mapeado na literatura, evidenciando uma mudança

profunda na rotina de investigação científica que ultrapassa uma simples correção textual. Diversos estudos têm se dedicado a categorizar esse fenômeno, como o estudo de Mishra2024, publicado no periódico *Scientific Reports*, no qual se investigou especificamente esse uso no contexto da pesquisa clínica, analisando as respostas de mais de 200 pesquisadores de 59 países. Os autores demonstraram que grande parte dos pesquisadores já possui familiaridade com LLMs, observando um impacto significativo dessas tecnologias no seu fazer científico. As pessoas autoras identificaram que os usos mais frequentes de LLM no fazer científico estão relacionados à revisão gramatical, à formatação e à edição estilística de textos. No entanto, há um uso crescente e relevante dessas ferramentas em tarefas cognitivas mais complexas, como a revisão de literatura, a análise de dados e a geração de novas ideias, nas quais as ferramentas de IA atuam como assistentes, aumentando a eficiência e economizando tempo no processo de pesquisa. Outros trabalhos expandem o uso de IA como ferramenta, como o de chen2025ai4research.

Nesse contexto, tomando como referência as categorizações propostas por Mishra2024 e chen2025ai4research, o presente trabalho propõe uma taxonomia própria para o uso de inteligência artificial ao longo do ciclo de vida da pesquisa científica, aderente também ao método hipotético-dedutivo de Popper, conforme apresentada na Figura 2.1.

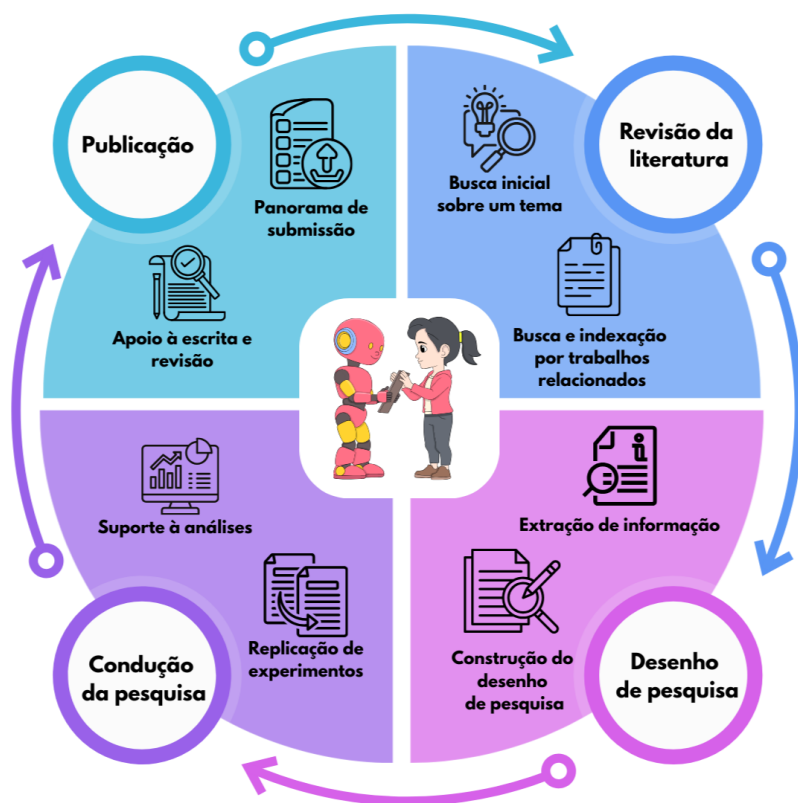


Figura 2.1: Taxonomia de tarefas de pesquisa utilizando a IA como um co-cientista, proposta neste trabalho.

A taxonomia apresentada na Figura 2.1 organiza esse ciclo em quatro categorias, a saber: i) IA para **Revisão da Literatura**, utilizada sobretudo para auxiliar na busca por trabalhos sobre o tema de interesse, além da indexação e recuperação semântica de traba-

lhos relacionados; ii) IA para **Desenho de pesquisa**, focado em ferramentas para extrair, interpretar e sintetizar conhecimento, e de auxílio na construção do desenho de pesquisa; iii) IA para a **Condução da pesquisa**, voltada para o suporte na replicação de experimentos e também na análise de interpretação de resultados; o ciclo finaliza com a iv) IA para **Publicação**, que engloba ferramentas para auxiliar na preparação, redação e finalização de manuscritos, oferecendo suporte desde a pré-revisão, meta-revisão e indicação de submissão. Vale destacar que essa taxonomia é uma adaptação dos modelos existentes, ajustada ao contexto de SI e adotada como base para a organização deste minicurso.

Essa estruturação está alinhada à evolução atual para a chamada IA agêntica, como apresentado no trabalho de gottweis2025towards. Neste estudo, os autores exploram uma arquitetura de sistema multiagente baseada no modelo Gemini 2.0, na qual a LLM atua de forma autônoma, desde o planejamento da pesquisa até a execução. O sistema opera por meio de um ciclo iterativo de raciocínio que tenta replicar o método científico, sendo: i) Agente de Geração - explora a literatura e propõe hipóteses; ii) Agente de Reflexão - tem o papel de revisor, verificando a novidade e a correção das hipóteses; iii) Agente de Ranking - organiza competições, nas quais as hipóteses passam por debates científicos simulados. A eficácia da proposta apresentada foi validada dentro dos grupos de pesquisa e também uma abordagem semelhante foi validada em cenários complexos por outros trabalhos, como o de [Gottweis et al. 2025a], que identificou de novos candidatos para o reposicionamento de fármacos no tratamento de um tipo de leucemia utilizando agentes, o que demonstra que a IA pode gerar hipóteses válidas em laboratório.

2.4. Panorama das ferramentas

A democratização do acesso ao ferramental foi um requisito norteador do minicurso, com o intuito de garantir que o conhecimento fosse acessível aos participantes. Dessa forma, as ferramentas selecionadas deveriam não apenas atender às funcionalidades requeridas para cada tarefa, mas também oferecer acesso gratuito ou permitir o uso de suas funções essenciais sem custos. A busca por ferramentas ocorreu a partir de um *insight* obtido durante a exploração de um artigo “kick-off”. A partir desse ponto, utilizou-se o Google Gemini, uma aplicação capaz de realizar pesquisas na internet e gerar relatórios contextualizados, incluindo as fontes que direcionam às ferramentas encontradas por meio da função chamada “Deep Research”. O *prompt* de entrada é apresentado na Figura 2.2

As soluções descritas no relatório retornado pelo Google Gemini incluem opções pagas e gratuitas, como observado na Figura 2.3. Foi então realizada uma curadoria para compreender o funcionamento, validar o uso e verificar a viabilidade dessas aplicações nas tarefas a serem exploradas. Dessa forma, foram selecionadas algumas delas, que serão descritas ao longo desta seção.

2.4.1. ChatGPT

O ChatGPT ¹ é uma aplicação de IA generativa que utiliza *deep learning* para a análise de contexto e a geração de conteúdo predominantemente textual. Projetado e treinado para produzir textos semelhantes aos elaborados por humanos, ele opera por meio de uma interface conversacional, na qual o usuário interage por meio de um texto, conhecido

¹<https://chat.openai.com>

Prompt para busca de ferramentas

Sou um **pesquisador iniciante em inteligência artificial** e, recentemente, tive contato com um artigo intitulado “Towards an AI co-scientist”. Fiquei interessado em outros trabalhos que envolvam IA como ferramenta auxiliadora do cientista, mais especificamente que me ajude em algumas atividades como:

- Busca, indexação e extração de informações de artigos científicos
- Construção de desenho de pesquisa, auxiliando na geração de hipóteses, perguntas de pesquisa até na definição de materiais e métodos
- Revisão de manuscrito, com sugestões de melhorias, locais de submissão e até atuar como um revisor para o periódico-alvo

O artigo que serviu de “kickoff” foi: Gottweis, J., Weng, W.H., Daryin, A., Tu, T., Palepu, A., Sirkovic, P., Myaskovsky, A., Weissenberger, F., Rong, K., Tanno, R. and Saab, K., 2025. Towards an AI co-scientist. arXiv preprint arXiv:2502.18864.

Por favor, me auxilie nesta busca, fornecendo:

- Sugestões de leituras
- Sugestões de ferramentas de IA
- Demais dicas.

Figura 2.2: *Prompt* de busca para ferramentas de auxílio a cientistas.

Ferramenta	Melhor Para	Tecnologia Chave	Custo (Estimado)
Elicit	Extração sistemática de dados, Matrizes de síntese.	Busca Semântica + LLM Extraction	Freemium / Assinatura para alto volume
SciSpace	Leitura interativa, explicação de conceitos, Rascunho.	Copilot for Science / ChatPDF	Freemium / Pro
ResearchRabbit	Visualização de redes de citação, Descoberta serendipituosa.	Graph Analysis	Gratuito
Consensus	Respostas rápidas baseadas em consenso científico ("Sim/Não").	NLP sobre Abstracts	Freemium
Paperguide	Relatórios de "Deep Research" automatizados.	Agentes Autônomos de Pesquisa	Assinatura

Figura 2.3: Captura de tela com as ferramentas resultantes do *prompt* da Figura 2.2.

como *prompt*, que é processado e respondido pelo modelo. O ChatGPT é de propriedade da OpenAI e, nesta edição do minicurso, foi utilizada a versão GPT-5, lançada em 07 de agosto de 2025. Uma das funcionalidades da ferramenta é o acesso a configurações especializadas do modelo base, desenvolvidas pela comunidade e fundamentadas em arquiteturas da família GPT. Entre essas configurações está o *Prompt Professor*², voltado

²<https://chatgpt.com/g/g-qfoOICq1l-prompt-professor>

à geração e à estruturação de *prompts* segundo os princípios da engenharia de *prompt*, compreendida como um conjunto de boas práticas que visa alinhar a saída do modelo ao resultado esperado.

Ressalta-se que tais customizações não configuram novos modelos no sentido estrito, uma vez que não envolvem modificações nos parâmetros nem na arquitetura subjacente, mas sim ajustes comportamentais induzidos por instruções e *templates* de interação. Essa abordagem é especialmente útil ao trabalhar com ecossistemas de LLMs, nos quais o *prompt* desempenha um papel central no processo. Apesar de disponibilizar uma série de funcionalidades ao usuário, o ChatGPT não é totalmente gratuito, adotando um caráter *freemium*. Convém pontuar que as limitações da versão gratuita não interferem no uso dessa configuração nas tarefas propostas neste minicurso.

2.4.2. Gemini

O Gemini³ é um modelo de LLM multimodal, pertencente ao Google. Ele opera por meio de uma interface de conversação para interagir com o usuário. A versão utilizada neste minicurso foi a Gemini 3 Pro, lançada em 26 de junho de 2025. Além da interação conversacional, essa versão possui a funcionalidade *Deep Research*, que permite estruturar uma *query* de busca para pesquisas na web sobre um determinado assunto, retornando um relatório com contexto e fontes.

Assim como outras ferramentas de IA generativa baseadas em chat, o Gemini utiliza o *prompt* como entrada, permitindo obter resultados mais concisos por meio da aplicação de técnicas de engenharia de *prompt*, como a definição de personas, a enumeração de etapas e a definição da formatação dos resultados. Em razão do modelo de negócio *freemium* adotado pelo Google, a versão 3 Pro, com todas as suas funcionalidades, não se encontra disponível no plano gratuito. Entretanto, o Google lançou o programa Google AI Pro⁴, no qual os estudantes podem obter um período de teste de um ano para utilizar as ferramentas, como o Google Gemini Pro. Além disso, sua natureza multimodal e a elevada capacidade de retenção de contexto tornam o Gemini uma opção particularmente adequada para buscas temáticas na web, o que fundamenta a solução adotada neste minicurso.

2.4.3. ResearchRabbit

A ResearchRabbit⁵ é uma plataforma de indexação e descoberta de artigos científicos, com um motor de busca voltado à expansão da literatura. Essa ferramenta possui uma função de coleção, na qual é possível definir um tema e inserir artigos relacionados ao domínio de interesse. Essas coleções podem ser compartilhadas, permitindo que grupos de pesquisa atuem de forma colaborativa. A partir de um trabalho-base, a plataforma apresenta um panorama, exibindo sugestões de novos artigos, os pesquisadores que citam aquele estudo em suas pesquisas e a opção de exportar as citações no formato BibTeX, para uso em trabalhos acadêmicos.

Adicionalmente, o ResearchRabbit gera essas relações de citação em formato de

³<https://gemini.google.com/app>

⁴<https://gemini.google/br/students/?hl=pt-BR>

⁵<https://app.researchrabbit.ai/>

grafo, no qual os nós representam os artigos e as arestas representam as citações entre eles. Esse grafo também pode ser organizado em formato de linha do tempo, o que facilita a compreensão da evolução da literatura. Dessa maneira, a aplicação permite uma navegação rápida pelas publicações e é especialmente útil para apoiar processos de *forward snowballing* e *backward snowballing*. Originalmente gratuita, a ferramenta passou a adotar o modelo *freemium*; ainda assim, suas funcionalidades essenciais permaneceram acessíveis a qualquer usuário que se cadastrar na plataforma.

2.4.4. NotebookLM

O NotebookLM⁶ é de propriedade do Google, que o classifica como um “assistente de pesquisa com tecnologia de IA”. Nele, é possível enviar arquivos de diversos tipos, como PDF. O NotebookLM os considera como única fonte e restringe suas respostas a esse conjunto de arquivos fornecidos, ajudando a reduzir alucinações. Essa ferramenta é especialmente interessante para a extração de informações, a sumarização de textos, a comparação entre trabalhos, dentre outras.

Uma aplicação de destaque é no estudo de um tema, ao realizar a síntese multimodal do conteúdo, produzindo vídeos, podcasts, *flashcards* e mapas mentais. Apesar de ser uma plataforma *freemium* e de possuir limitações de funcionalidades no plano básico, a escolha se deu por conta da integração com o Google e com o Google AI Pro, que possibilitam o acesso *premium* com licença de estudante.

2.4.5. DeepSeek

O DeepSeek⁷ é um modelo de linguagem de grande porte, de propriedade da empresa chinesa DeepSeek Company, atualmente na versão DeepSeek-V3.2. Essa ferramenta oferece uma interface conversacional que permite ao usuário interagir diretamente com o modelo por meio de *prompts*, de forma semelhante a outras LLMs. Entre seus diferenciais, está a geração de respostas estritamente textuais, com foco acentuado no raciocínio técnico, o que lhe confere destaque em tarefas como a programação [Deng et al. 2025].

A escolha desse modelo se deu principalmente por sua aptidão para o desenvolvimento de código e por ser gratuito. Tendo sua limitação atrelada a *tokens* de aproximadamente 128 mil por sessão, esse fato pode ser facilmente contornado por meio da abertura de novos chats para a continuidade da construção de artefatos de software.

2.4.6. Perplexity AI

O Perplexity AI⁸ funciona como um mecanismo de busca baseado em inteligência artificial; a plataforma integra LLMs como Claude Sonnet 4.5, Grok 4 e GPT-5, permitindo que o usuário selecione o modelo mais adequado à tarefa. Para gerar respostas, utiliza técnicas de recuperação de informação, de sumarização e de análise contextual, combinando a consulta à web com o processamento de linguagem natural. Entre suas funcionalidades está o Pro Search, que amplia o conjunto de fontes consultadas e apresenta um relatório estruturado com referências, funcionalidade semelhante à de outros agentes de pesquisa

⁶<https://notebooklm.google.com/>

⁷<https://chat.deepseek.com/>

⁸<https://www.perplexity.ai/>

baseados na web.

O Perplexity AI possui natureza multimodal, o que permite trabalhar com documentos em PDF, imagens, textos etc. A solução adota o modelo *freemium*, oferecendo um conjunto básico de funcionalidades gratuitas. Esta ferramenta também conta com um programa para estudantes, que concede acesso ao plano *premium* por um período de teste de 1 ano. A integração com múltiplos LLMs contribui para um bom desempenho em tarefas como a geração de código, a síntese de literatura, a elaboração de relatórios e o apoio a atividades científicas em geral. O Perplexity AI se mostra uma opção versátil no ecossistema de ferramentas de apoio à pesquisa, complementando as demais apresentadas nesta seção. Conforme sintetizado na Tabela 2.1, observa-se um conjunto diverso de recursos que auxiliam o cientista em diferentes etapas do processo de pesquisa científica.

Tabela 2.1: Síntese do Panorama das ferramentas utilizadas no minicurso.

Ferramenta	Função Principal	Vantagens	Limitações
ChatGPT	Apoio na engenharia de <i>prompt</i> .	Modelos personalizados especializados em tarefas específicas.	Algumas limitações na versão gratuita e recursos avançados restritos.
Gemini	Busca e indexação de literatura com relatórios e fontes.	Integração com o Google e acesso gratuito para estudantes.	Necessita de uma conta do Google; acesso premium após o período de teste.
DeepSeek	Replicação de experimentos e de tarefas de programação.	Gratuito, com bom desempenho técnico e lógico.	Limite de <i>tokens</i> .
Perplexity AI	Busca na literatura e no auxílio em tarefas de programação.	Versatilidade: permite a seleção entre diversos modelos.	Recursos e modelos mais avançados estão disponíveis apenas na conta <i>premium</i> .
ResearchRabbit	Expansão e organização de coleções de artigos científicos.	Gratuito com opção de colaboração e exportação em formato BibTeX.	Tornou-se <i>Freemium</i> e alguns recursos avançados são pagos.
NotebookLM	Extração e síntese das informações de artigos enviados.	Reduz alucinações, gera resumos e mapas mentais e integra-se ao Google.	Algumas funções restritas à versão <i>premium</i> requerem uma conta do Google.

Os exemplos de uso das ferramentas sintetizadas na Tabela 2.1 são apresentados na Seção a seguir.

2.5. Utilizando as ferramentas como um co-cientista

Nesta Seção, serão apresentados exemplos de uso considerando os quadrantes propostos na Figura 2.1. Com base no levantamento das principais plataformas e de suas funcionalidades, definiram-se seis tarefas. As duas primeiras estão relacionadas ao primeiro quadrante (**Revisão da Literatura**) e abordam, respectivamente, a imersão em um tema com o Gemini e a indexação e expansão da literatura com o ResearchRabbit. Partindo para o segundo quadrante, **Desenho de Pesquisa**, abordam-se as tarefas de extração de informações de artigos com o NotebookLM e de construção de desenhos de pesquisa com o Gemini. Na etapa de **Condução da Pesquisa**, apresenta-se a reprodução de experimentos com o DeepSeek, o GPT e o Gemini. Por fim, no quadrante de **Publicação**,

abordam-se a revisão de artigos e o suporte no processo de submissão no Gemini.

2.5.1. Imergindo em um tema

Conforme mencionado, a primeira tarefa consiste em utilizar o Gemini para aprofundar o tema abordado no trabalho base. Ao utilizar *prompts* específicos, é possível delimitar o escopo da revisão bibliográfica, identificar ferramentas relevantes, encontrar termos de busca adequados e localizar autores atuantes na área de IA aplicada à pesquisa acadêmica. O primeiro exemplo de *prompt* apresentado neste capítulo, especificamente na Figura 2.2, ilustra o uso da funcionalidade *Deep Research* do Gemini para buscar informações pertinentes na internet e organizar o conhecimento de forma estruturada em um relatório. O *prompt* inicia com a definição de uma persona, no caso, “pesquisador iniciante em inteligência artificial”, e passa o artigo-base, solicitando outros trabalhos e ferramentas que envolvam IA como auxiliar do cientista. Ainda, o *prompt* apresentado na 2.2 lista algumas tarefas, como busca, indexação e extração de informações de artigos, construção de um desenho de pesquisa e até revisão de manuscritos.

A funcionalidade *Deep Research* leva alguns minutos, considerando seu perfil agêntico de busca na web. A resposta da LLM contextualiza o tema e apresenta uma série de hiperlinks que levam ao local de onde essas informações foram extraídas. O relatório gerado pelo *Deep Research* do Gemini pode ser visto na Figura 2.4.



Figura 2.4: Relatório gerado com o *Deep Research* do Gemini.

Conforme visto na Figura 2.4, esse tipo de relatório auxilia em diversas etapas da pesquisa acadêmica, especialmente na revisão da literatura e na identificação de trabalhos correlatos. Ele facilita a busca e o levantamento de publicações relevantes, organizando essas informações de forma estruturada. Isso permite ao pesquisador focar diretamente

nos trabalhos pertinentes ao tema investigado. Adicionalmente, a ferramenta sumariza todo o conteúdo gerado em um quadro, como se observa na Figura 2.3 (apresentada na Seção de ferramentas). Essa representação ajuda na seleção de trabalhos de forma mais otimizada, pois permite uma avaliação rápida das bases de dados, das ferramentas e do contexto de aplicação utilizados.

2.5.2. Expansão da literatura

Em continuidade à imersão no tema, o passo seguinte é a expansão da literatura; para isso, será necessário criar uma coleção na ferramenta ResearchRabbit, a fim de organizar os trabalhos indicados pelo Gemini a partir do Relatório e da Sumarização (Figuras 2.4 e 2.3). Utilizando os recursos do ResearchRabbit, é possível visualizar artigos semelhantes, as redes de citações e os autores relacionados. Essas informações facilitam a exploração da literatura, a atualização das citações, a identificação de publicações relevantes e de trabalhos seminais. Na Figura 2.5, observa-se uma coleção de artigos extraídos das indicações do relatório elaborado no passo anterior.

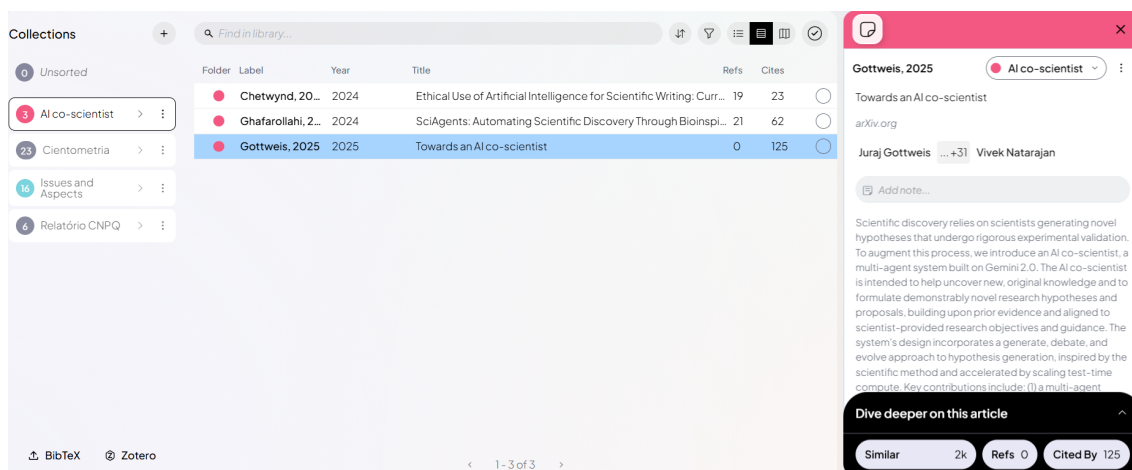


Figura 2.5: Coleção de artigos organizada no ResearchRabbit.

Conforme visto na Figura 2.5, a coleção é composta por 3 trabalhos. Há cinco colunas principais: o identificador (“*Label*”) de cada um, o ano de publicação (“*Year*”), o título do trabalho (“*Title*”), o número de referências que a ferramenta conseguiu identificar (“*Refs*”) e, por fim, o número de citações (“*Cites*”). Estes dois últimos atributos são os principais utilizados pela ferramenta para indicar artigos semelhantes. Esta é uma outra funcionalidade. Ao navegar pelos artigos similares, é possível visualizar a rede de citações e os artigos que o citaram, como apresentado na Figura 2.6.

Conforme pode ser observado na Figura 2.6, os trabalhos são listados à esquerda e a rede de interação entre os artigos é representada pelo grafo à direita, que pode ser manipulado para melhor visualização. Com os dados apresentados pela plataforma, o pesquisador pode buscar rapidamente na literatura e observar o panorama dos trabalhos sobre o tema. Essa navegação rápida otimiza o processo de busca e de indexação da literatura. Adicionalmente, a plataforma possibilita a integração com o Zotero, um sistema de indexação de artigos, e também permite a exportação das citações e referências em diversos formatos, inclusive para LaTeX/BibTeX, o que simplifica a criação da seção de

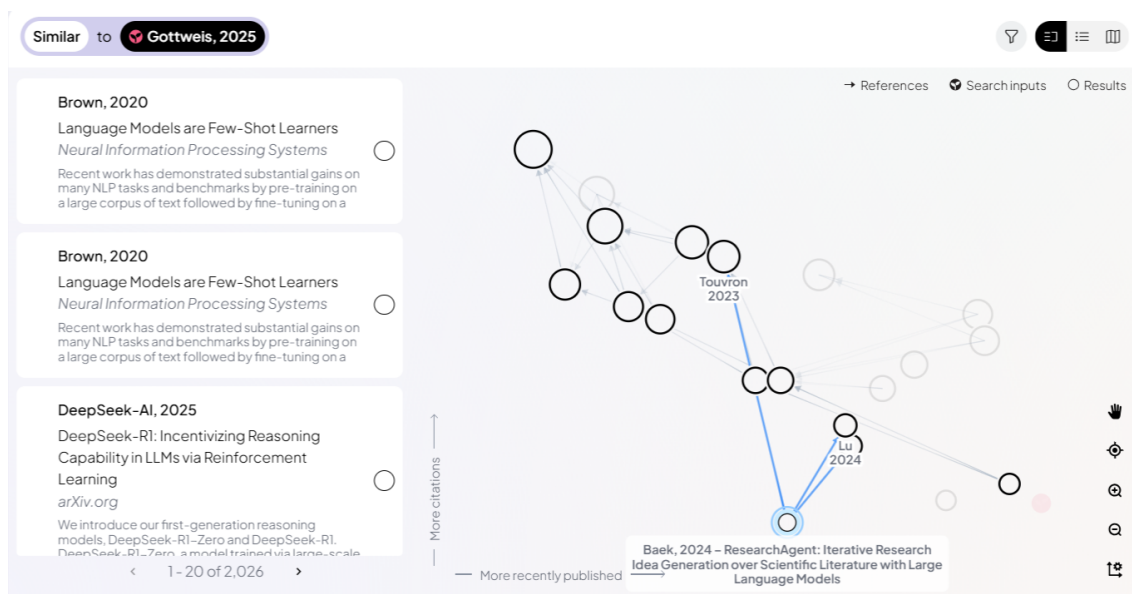


Figura 2.6: Artigos similares indicados pelo ResearchRabbit.

referências. Após ampliar a biblioteca de artigos com o relatório do Gemini e as sugestões de trabalhos do ResearchRabbit, o próximo passo é baixar esses arquivos para proceder com a extração de informações.

2.5.3. Extração de informações

A extração de informações é realizada com o apoio do NotebookLM. Tomando como base o artigo inicial, as recomendações do Gemini *Deep Research*, a expansão com o ResearchRabbit e a curadoria pelas pessoas pesquisadoras, um total de cinco artigos representativos foi selecionado para a extração de informações, conforme mostrado na Figura 2.7.

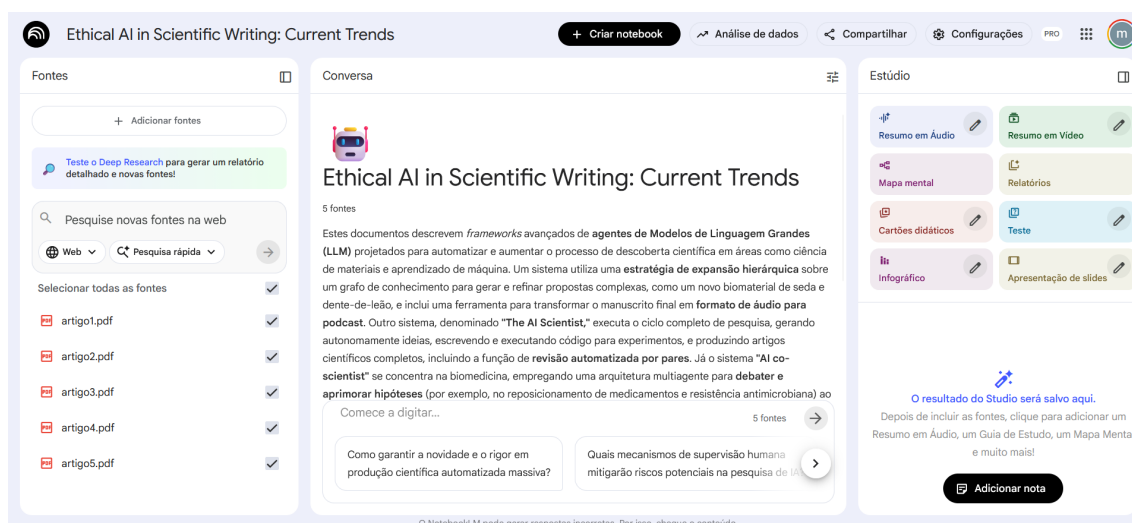


Figura 2.7: Tela inicial do NotebookLM.

A Figura 2.7 apresenta a tela após o *upload* dos artigos mencionados. Na área à

esquerda (Fontes), os artigos são listados. A área central apresenta a interface de *chat*, com um breve resumo do tema. A área à direita apresenta o Estúdio, com funcionalidades para resumo em áudio, vídeo, mapa mental ou cartões didáticos, apenas para citar alguns. Estas funcionalidades auxiliam na compreensão de conceitos ou de novas modalidades de compreensão. Também pode ser observado o formato multimodal das saídas possíveis, além da caixa de texto onde podem ser enviadas perguntas acerca dos arquivos-fonte, como, por exemplo: “Como os sistemas de agentes de IA aprimoram a criatividade e o rigor na geração de hipóteses científicas?”. A resposta a essa requisição está representada na Figura 2.8.

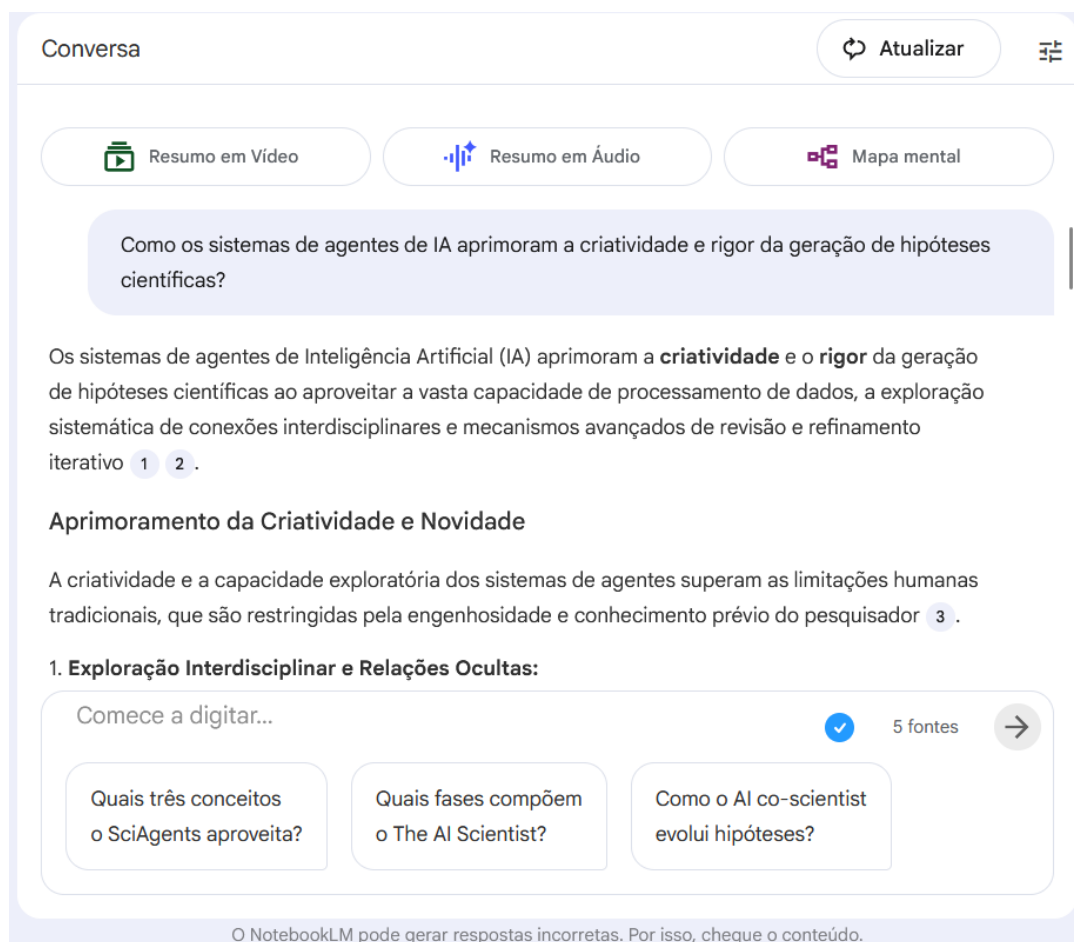


Figura 2.8: Resposta a uma requisição no NotebookLM.

Conforme se pode ver na Figura 2.8, a resposta à pergunta é acompanhada de números em cinza, como no primeiro parágrafo. Eles são hiperlinks para trechos do(s) artigo(s) e funcionam como um RAG. Ainda, na resposta apresentada na Figura 2.8, é possível observar que o NotebookLM exibe um aviso, na parte inferior da imagem, informando que a ferramenta pode cometer erros. Por este motivo, é de grande importância a presença humana nessas tarefas, visando a curadoria das respostas geradas pelos modelos. Dessa maneira, o NotebookLM demonstra-se uma ferramenta interessante para a extração de informações de artigos. Ressalta-se que as ferramentas apresentadas ao longo do minicurso auxiliam o cientista e otimizam suas tarefas, mas, em nenhum momento,

devem ser usadas sem supervisão e verificação de seus resultados.

2.5.4. Desenho de pesquisa

Além do cenário de busca a partir de um artigo *kick-off*, as LLMs podem auxiliar em outras tarefas, como na definição ou no aprimoramento do desenho de pesquisa. Aqui utilizaremos um exemplo para auxiliar na construção de um instrumento de coleta de dados. A Figura 2.9 apresenta este *prompt*, enviado ao Gemini na versão 3 Pro, na configuração padrão, sem uso de ferramentas adicionais.

Prompt para auxiliar na construção do desenho de pesquisa survey

Você foi contratado para **ministrar um curso sobre o uso de IA na Gestão Pública**, com enfoque em **Gestão Acadêmica e de Hospitais Universitários**. Foi solicitado que houvesse duas dinâmicas: a primeira é de caráter diagnóstico; a segunda, de caráter propositivo.

No diagnóstico, a ideia é aplicar um **questionário autorespondido**, com tempo de resposta estimado em 3 minutos, coletando variáveis como: com qual **frequência os participantes utilizam ferramentas de IA**, para quais **tarefas** (com uma lista de tarefas que as pessoas mais usam), quais as **ferramentas** (com a lista das ferramentas mais usadas), como os participantes **validam/usam resultados** (lista de formas), como elas **avaliam** a efetividade das respostas e outros aspectos que julgares pertinentes.

Neste momento, apresente o questionário para que seja estruturado em uma ferramenta adequada.

Figura 2.9: Prompt para auxílio na construção do desenho de pesquisa baseada em *survey*.

O *prompt* apresentado na Figura 2.9 segue a linha estabelecida anteriormente sobre as boas práticas que levam a uma resposta mais assertiva, e espera-se, como resultado, uma lista para realizar uma pesquisa do tipo *survey* para aquisição de informações. A taxonomia do *prompt* é composta por definição da persona, solicitação e contexto dos artefatos solicitados. O resultado da solicitação pode ser observado na Figura 2.10. Também é possível observar o mesmo aviso no NotebookLM sobre o erro associado ao uso de LLMs, reforçando ainda mais a necessidade de verificação humana das respostas geradas. Com isso, nota-se que é possível utilizar LLMs em mais tarefas, incluindo as relacionadas ao desenho da pesquisa.

Conforme se observa na Figura 2.10, o Gemini fornece claramente o instrumento de coleta de dados, organizado em seções e com os tipos de perguntas bem definidos. Destaca-se que este mesmo *prompt* foi utilizado em um cenário real, em que as pessoas autoras iriam dar uma palestra sobre o tema e realizar uma coleta de dados para caracterizar o público-alvo. Para melhorar a confiabilidade, o *prompt* também foi submetido a outras ferramentas, o que permitiu uma diversidade de respostas, cabendo às pessoas autoras a curadoria delas e a elaboração do questionário, tomando como ponto de partida os resultados das LLMs.

Além de auxiliar no processo de desenho da pesquisa quando se trata de instrumentos de coleta de dados, como *surveys*, as LLMs também podem auxiliar o pesquisador no desenvolvimento de perguntas de pesquisa, que, em geral, são o ponto de partida ao iniciar um projeto de investigação. Para esse fim, o usuário pode elaborar um *prompt* que forneça contexto e domínio, com base nas buscas realizadas no Gemini *Deep Research*,

Estrutura do Questionário Diagnóstico

Título Sugerido: Mapeamento de Maturidade em IA - Gestão Acadêmica e Hospitalar **Descrição:**

Este questionário é anônimo e visa personalizar o conteúdo do nosso curso às suas necessidades reais. Não existem respostas certas ou erradas.

Seção 1: Perfil Básico

(Para segmentar se as dores são administrativas ou clínicas/acadêmicas)

1. Em qual área você atua predominantemente?

- ☐ Gestão Administrativa/Financeira
- ☐ Gestão Acadêmica (Secretaria, Coordenação, Pós-Graduação)
- ☐ Assistencial/Hospitalar (Gestão de Leitos, Atendimento, Farmácia)
- ☐ Tecnologia da Informação / Dados
- ☐ Outro

Seção 2: Frequência e Ferramentas

(Coleta das variáveis de hábito solicitadas)

2. Com qual frequência você utiliza ferramentas de Inteligência Artificial Generativa (como

Peça ao Gemini

+ Ferramentas Raciocínio

O Gemini pode cometer erros. Por isso, é bom checar as respostas.

Figura 2.10: Questionário resultante da requisição ao Gemini.

na expansão por meio do ResearchRabbit e na extração de conceitos e terminologia com o NotebookLM. Um exemplo da aplicação da requisição de formular perguntas de pesquisa pode ser observado na Figura 2.11.

Prompt para auxiliar na elaboração de perguntas de pesquisa

Você deverá atuar como **especialista em Inteligência Artificial**, com sólida **experiência na elaboração de projetos de pesquisa** voltados à captação de recursos. Tendo em vista que a governança da IA é um campo interdisciplinar que busca orientar o desenvolvimento e o uso da inteligência artificial de forma segura, ética, justa e benéfica para a humanidade: Apresente possíveis **perguntas de pesquisa sobre Governança de IA** para problemas de hoje e do futuro.

Forneça perguntas de pesquisa instigantes e relevantes para a academia. Apresente uma breve explicação, mencionando as perspectivas teóricas e os métodos de pesquisa adequados para respondê-las.

Figura 2.11: *Prompt* para auxiliar na elaboração de uma pergunta de pesquisa.

A taxonomia de *prompt* adotada começa pela definição da persona, do contexto e da requisição. O resultado gerado a partir dessa estrutura é parcialmente apresentado na Figura 2.12.

💡 Perguntas de Pesquisa Estratégicas em Governança de IA

1. Eixo: Regulação e Vieses Algorítmicos (Problema Atual)

Pergunta de Pesquisa:

"Como os frameworks regulatórios emergentes (ex: 'AI Act' da União Europeia ou o PL brasileiro) podem efetivamente garantir a **auditabilidade** e a **mitigação de vieses** em modelos de Machine Learning de uso geral (foundation models) que operam como caixas-pretas opacas?"

- **Relevância:** A dificuldade de fiscalizar modelos massivos e proprietários é o maior desafio regulatório de hoje. A legislação exige transparência e justiça, mas os métodos técnicos de **Explainable AI (XAI)** ainda são limitados.
- **Perspectivas Teóricas:**
 - **Ética da Informação e Justiça Distributiva:** Foco na equidade de resultados do sistema.
 - **Teoria da Regulação:** Análise da efetividade de instrumentos *ex ante* (antes do uso) versus *ex post* (após o uso).
- **Métodos Sugeridos:**
 - **Auditoria Técnica de Modelos (Adversarial Testing):** Desenvolvimento de *benchmarks* para identificar e quantificar o **viés algorítmico** em modelos de linguagem para o contexto local.
 - **Análise de Discurso e Jurimetria:** Comparação da linguagem legal com a capacidade técnica de *compliance*.

5. Eixo: Cidadania Algorítmica e Inclusão (Problema Atual)

Pergunta de Pesquisa:

"De que forma os sistemas de **Decisão Algorítmica (ADS)** usados na gestão pública (saúde, educação, segurança) reconfiguram a noção de **Cidadania**, e quais são os requisitos de **transparência radical** e **direitos de recurso humano** necessários para preservar a agência dos cidadãos?"

- **Relevância:** A IA é cada vez mais usada em decisões críticas do Estado (concessão de benefícios sociais, triagem médica, vigilância). A falta de transparência cria uma "cidadania de segunda classe" para os indivíduos afetados pelas decisões da caixa-preta.
- **Perspectivas Teóricas:**
 - **Teoria da Democracia e Esfera Pública:** O acesso à informação e a capacidade de contestação.
 - **Sociologia do Risco:** A distribuição desigual dos riscos impostos pela tecnologia.
- **Métodos Sugeridos:**
 - **Estudos de Caso (Case Studies):** Análise detalhada de sistemas de ADS específicos em cidades ou estados, focando na experiência dos usuários finais.
 - **Pesquisa-Ação Participativa:** Envolver grupos de cidadãos afetados na cocriação de mecanismos de *feedback* e recurso contra decisões algorítmicas.

Figura 2.12: Resposta Gemini pergunta de pesquisa.

A análise do resultado apresentado na Figura 2.12 evidencia que as perguntas são organizadas em eixos temáticos, acompanhadas da indicação de relevância, de perspectivas teóricas e de métodos sugeridos. Esta funcionalidade é bastante útil para a área visando o enquadramento da teoria em SI. Por exemplo, o SBSI aponta para a lista de Teorias de SI disponível no portal da BYU Library⁹, que, no momento dessa consulta, apresenta 89 teorias. É possível também incluir a lista no *prompt*. Reitera-se que a *expertise* humana no julgamento de pertinência é imprescindível.

Ainda no desenho de pesquisa, é comum a presença de ameaças à validade, decorrentes de fatores não considerados ou de pontos influentes capazes de afetar os resultados. Caso não sejam devidamente identificadas e mitigadas, essas ameaças podem comprometer a robustez e a credibilidade dos achados. Nesse contexto, as LLMs podem apoiar o pesquisador na identificação de ameaças à validade. Com o uso da ferramenta Gemini, é possível realizar consultas para identificar tais pontos, enviando o arquivo da pesquisa acompanhado do *prompt* apresentado na Figura 2.13.

Prompt para identificar pontos que possam representar ameaças à validade da pesquisa.

Preciso identificar as **ameaças potenciais à validade do meu trabalho**. Vou enviar um exemplo de trabalho que utiliza essa estrutura para construir uma seção. O exemplo se intitula: Detecção e priorização de problemas com base na análise de aplicativos móveis. E aqui está o meu trabalho.

Figura 2.13: *Prompt* para auxílio na elaboração de perguntas de pesquisa.

⁹<https://guides.lib.byu.edu/>

O resultado desta solicitação consiste numa lista detalhada dos possíveis pontos que representam ameaças à validade da pesquisa apresentada. Interessante notar que a resposta a essa solicitação foi gerada em inglês, por ser o idioma do material enviado. O resultado é apresentado na Figura 2.14.

Based on my analysis of your work, "Fintechs' issues," and using the framework from Wohlin et al. (2012) as demonstrated in the example paper, here are the potential threats to the validity of your study.

Threats to Validity

In this section, we discuss the potential threats to the validity of our study. We follow the guidelines proposed by Wohlin et al. (2012), categorizing threats into four types: construct validity, internal validity, external validity, and conclusion validity.

Internal Validity

Internal validity is concerned with the causal relationship between the treatment and the outcome, questioning whether our methodology was the true cause of the results, or if other factors influenced them.

- **Instrumentation:** The primary instruments in this study are the LLMs themselves (Gemma, LLAMA, GPT, DeepSeek). Proprietary, very-large-size models are particularly problematic as they are "black boxes" and can be updated by their providers without notice. This means that running the same prompt at different times could yield different results, affecting the replicability and stability of the findings. This was partially mitigated by setting the temperature parameter to zero where possible to increase determinism.
- **Human-in-the-Loop Bias:** The heavy reliance on domain experts for guidance, validation, and interpretation is a core feature of the methodology but also an internal threat. The experts' subjective judgments, prior experiences, and potential biases could be a confounding variable that influences which issues are ultimately considered valid and how they are grouped. The final list of issues is a product of both the LLM's processing and the experts' cognition.

External Validity

External validity concerns the generalizability of the experimental results to other contexts.

- **Interaction of Selection and Treatment:** The study is focused on a specific context: three selected Fintechs within the Brazilian market. The types of issues, their frequency, and their relationship to specific products may not be generalizable to:
 - Other Fintechs in Brazil (e.g., those with different business models).
 - Traditional incumbent banks.
 - Fintechs in other countries with different financial regulations, products (e.g., PIX is specific to Brazil), and cultural contexts for complaining.
- **Data Source Specificity:** As the data is sourced exclusively from `consumidor.gov.br`, the findings may only be generalizable to issues reported on similar, government-run, formal dispute resolution platforms. The framework's performance and the types of issues discovered might differ significantly if applied to less structured data from social media or app reviews.

Conclusion Validity

Conclusion validity refers to the ability to draw correct conclusions about the relationship between the treatment and the outcome.

- **Lack of Statistical Measures:** The evaluation of the framework is primarily qualitative, relying on expert analysis, clustering, and comparison to an official list from the Brazilian Central Bank. The study does not employ quantitative statistical tests to measure the "accuracy" or "effectiveness" of the issue extraction process. Therefore, conclusions about the superiority of one prompt or clustering method over another are based on descriptive and interpretive analysis rather than statistical significance.

Figura 2.14: Resposta do Gemini referente ao *prompt* para identificar ameaças à validade da pesquisa.

Conforme observado na Figura 2.14, a LLM sugeriu potenciais ameaças à validade do estudo de acordo com as categorias propostas pelo trabalho-referência. É importante notar que, apesar de serem bem escritas e convincentes, após uma discussão interna, foi perceptível que diversas ameaças sugeridas não eram aplicáveis e outras foram negligenciadas. De toda sorte, a ferramenta mostrou-se interessante para a ideação e para o auxílio na construção do desenho de pesquisa, encerrando o segundo quadrante da taxonomia proposta na Figura 2.1.

2.5.5. Replicação de experimentos

Passando ao quadrante de Condução da Pesquisa (Figura 2.1), apresenta-se a parte relacionada à replicação de experimentos com LLMs. Considerando a complexidade inerente a este processo, esta atividade é dividida em três partes. A primeira consiste em extrair os passos metodológicos adotados no trabalho que se deseja reproduzir. A segunda etapa consiste em transformar os passos em *prompts* estruturados, contendo todos os requisitos lógicos e técnicos necessários. O terceiro passo é a construção dos códigos usando como instrução os *prompts* criados no passo anterior.

Para a extração das etapas, foi utilizado o ChatGPT, em conjunto com a ferramenta Professor-Prompt, para auxiliar no refinamento dos comandos a serem passados à LLM. Esse *prompt* será responsável pela análise e pela extração das informações metodológicas do estudo original. Essa estruturação decorre da necessidade de uma resposta em um

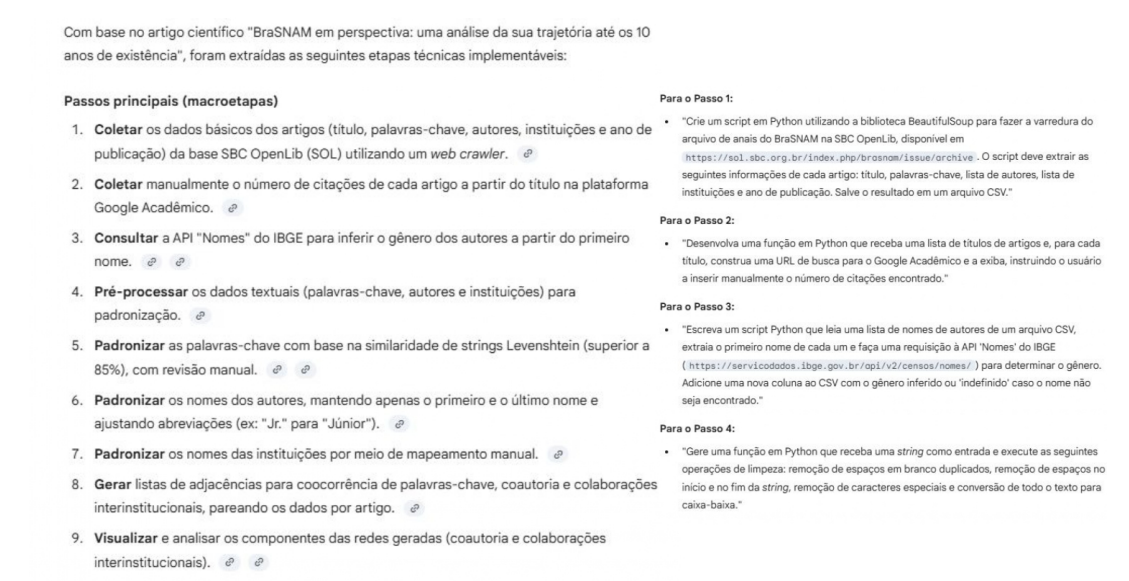


Figura 2.15: Resultado da extração de etapas metodológicas.

formato específico. Após a criação do *prompt* estruturado, o próximo passo é enviar ao Gemini juntamente com o arquivo PDF do artigo que se deseja replicar. A partir desse processo, extraem-se as tarefas descritas na seção de metodologia do trabalho original, conforme ilustrado na Figura 2.15.

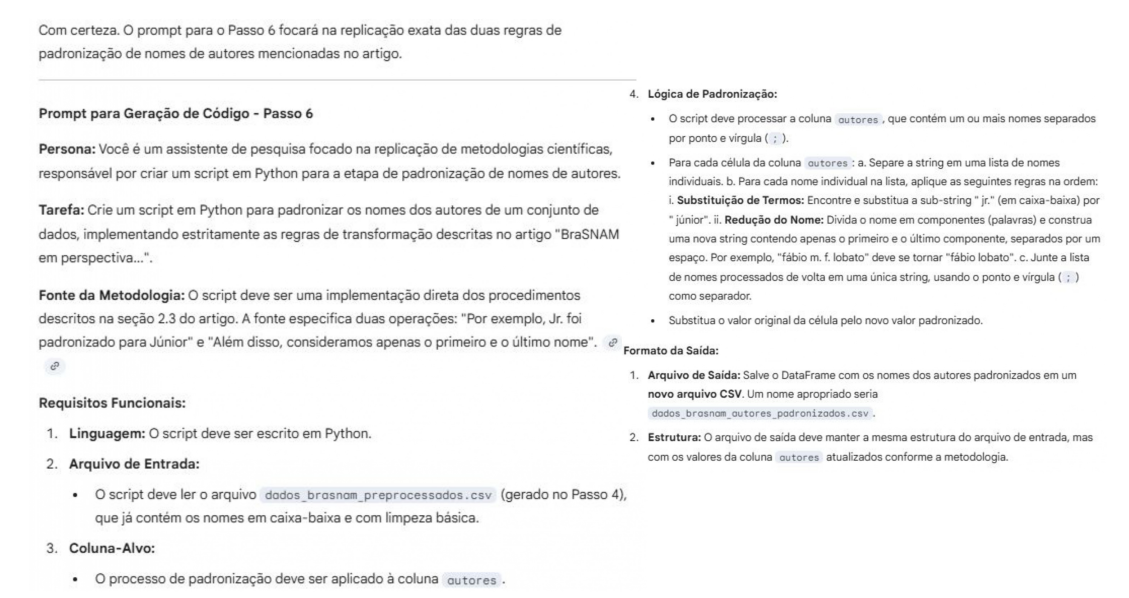


Figura 2.16: Geração de *prompts* específicos com Gemini.

Na segunda etapa foi feita uma requisição ao Gemini para gerar individualmente os *prompts* para cada uma das etapas extraídas no passo anterior. A Figura 2.16 apresenta um exemplo de um dos passos extraídos do trabalho original. Basicamente, refere-se a um processo de padronização de dados, no qual são descritos a linguagem de programação

desejada, o tipo de dado com que se quer trabalhar, nome do arquivo e toda a parte lógica que deve ser implementada nesse código. Assim, é possível utilizar a LLM para extrair o fluxo experimental do trabalho base e, a partir disso, gerar instruções estruturadas para que outra LLM gere os códigos necessários à execução dos experimentos.

Na terceira etapa, o resultado das requisições realizadas nas fases anteriores é utilizado para instruir o DeepSeek, que fica encarregado da geração dos códigos necessários para reproduzir os passos técnicos descritos no trabalho-fonte. Neste estudo, foi utilizado o trabalho de lobato2021brasnam para garantir a replicabilidade dos experimentos, uma vez que o estudo em questão é de autoria do grupo de pesquisa. Os *prompts* gerados pelo Gemini foram inseridos no DeepSeek, com a funcionalidade de raciocínio ativada, com o intuito de gerar códigos mais precisos. Esse processo é apresentado na Figura 2.17, na qual o código é construído conforme as instruções do *prompt*, incluindo bibliotecas, linguagens de programação e lógica.

Passo 6 - Script Python para padronizar nomes de autores

Aqui está o script Python que implementa estritamente a metodologia descrita no artigo:

```
python

import pandas as pd

# Carregar o arquivo de entrada
df = pd.read_csv('dados_brasnam_preprocessados.csv')

# Função para padronizar um único nome de autor
def padronizar_nome(nome):
    # 1. Substituir " jr." por " júnior"
    nome = nome.replace(' jr.', ' júnior')

    # 2. Reduzir para primeiro e último componente
    componentes = nome.split()
    if len(componentes) >= 2:
        nome_reduzido = componentes[0] + ' ' + componentes[-1]
    else:
        nome_reduzido = componentes[0] if componentes else ''
```

Mensagem para DeepSeek

☒ Pensamento Profundo ☐ Pesquisar

Gerado por IA, apenas para referência

Figura 2.17: Geração de código para os passos metodológicos com DeepSeek.

Apesar de todo o cuidado na elaboração das instruções fornecidas ao modelo, notou-se que, durante a execução deste passo com o artigo de exemplo, em alguns pontos o código gerado apresentou falhas no fluxo de execução, o que demandou uma intervenção humana. Para tanto, foi imprescindível a compreensão integral do fluxo metodológico do trabalho, a fim de solicitar os ajustes necessários. Após a realização dessas intervenções, o

código passou a executar conforme o esperado, produzindo resultados satisfatórios. Dessa forma, evidencia-se o potencial do uso de LLMs como ferramentas de apoio em diversos aspectos do cotidiano do cientista, ainda que sem a plena automatização e com a necessidade de supervisão e intervenção humanas ao longo de todo o processo. Como o suporte às análises e interpretações dos resultados é semelhante à replicação de experimentos, a seção a seguir passa para o último quadrante do fluxo proposto na Figura 2.1.

2.5.6. Publicação

Um dos usos mais disseminados das LLMs na academia é, de forma geral, a automação de atividades relacionadas à escrita científica. Estendemos o suporte à escrita e à revisão, incluindo também o apoio à submissão de artigos. Considerando as particularidades inerentes ao processo de escrita e ao pleno uso, focaremos na revisão de textos e no auxílio à submissão. O *prompt* foi elaborado seguindo o mesmo padrão adotado nas etapas anteriores. Esse *prompt* inclui a definição de persona, o escopo e a apresentação da tarefa geral e das subtarefas, e termina com a saída esperada. Dessa forma, todo o raciocínio por trás da tarefa é organizado, como se observa na Figura 2.18.

Prompt para auxiliar na revisão escrita do trabalho.

Persona: Você é um revisor acadêmico experiente, com atuação em periódicos científicos de prestígio, especializado em cientometria, bibliometria e engenharia elétrica.

Contexto: Finalizamos um artigo de cientometria e bibliometria sobre pesquisadores com bolsa de produtividade que atuam na área de engenharia elétrica. O artigo foi escrito inicialmente em português e depois traduzido para o inglês; dessa forma, algumas falhas terminológicas e inconsistências podem ter sido introduzidas no processo de tradução.

Tarefa Geral: Você deverá atuar como revisor e elaborar três relatórios distintos, conforme descrito a seguir.

Relatório 1 – Análise Terminológica Crie um quadro contendo: eventuais falhas terminológicas; Sugestões de termos mais utilizados na literatura acadêmica; Uma justificativa para cada uma das falhas e sugestões.

Relatório 2 – Clareza e Fluidez Textual Selecione trechos de frases que não soam bem ou estão confusos. Sugira melhorias, apresentando uma pequena explicação.

Relatório 3 – Revisão por Pares (*Peer Review*) Atue como revisor de periódico e realize uma revisão criteriosa, considerando os critérios acadêmicos de periódicos prestigiosos.

Figura 2.18: *Prompt* para auxílio na revisão da redação do trabalho.

O resultado obtido a partir do *prompt* apresentado na Figura 2.18 consiste em uma lista de pontos de melhoria, com justificativas e sugestões de novos trechos. Assim, torna-se possível aprimorar o próprio processo de escrita. Tal estratégia é importante para evitar que a pessoa pesquisadora fique dependente da tecnologia e também para evitar vícios de linguagem característicos das LLMs, garantindo a exclusividade da autoria do texto a um humano. Outrossim, evitam-se outros problemas da escrita inteiramente artificial, como a possibilidade de plágio não intencional, deturpação argumentativa, entre outros.

Outro uso bastante interessante é no auxílio à submissão. Aqui, abordamos duas tarefas críticas: a escolha do local de publicação e a simulação de revisão por pares. O *prompt* inicia com a definição da persona (Pesquisador com experiência internacional), segue com a definição da tarefa, com todas as orientações de passos a serem seguidos pelo

modelo e, por fim, delimita-se o tipo de saída esperada. O *prompt* completo é ilustrado na Figura 2.19.

Prompt para sugestão de locais de publicação.

Persona: Você é um pesquisador experiente em publicações científicas nas áreas de inteligência artificial, linguagem natural, com amplo conhecimento sobre periódicos relevantes internacionalmente e no Brasil.

Contexto: Desenvolvi um artigo cujo título e resumo são: *Title:* Ferramentas de Inteligência Artificial para suporte ao processo de pesquisa em sistemas de informação *Abstract:* (resumo do artigo fornecido)

O objetivo é identificar periódicos adequados para a submissão desse artigo, considerando impacto acadêmico, visibilidade e alinhamento temático.

Tarefa: Elencar periódicos científicos interessantes para submissão do artigo, organizando-os em uma tabela com informações bibliográficas e razões de escolha. **Periódicos Seleccionados:** Para cada periódico listado na tabela a seguir, extraia ou estime os seguintes itens: **Nome do periódico;** **Fator de impacto** (ou, se disponível, índices bibliométricos equivalentes, como SJR, JIF ou outro indicador relevante); **Estrato Qualis (CAPES)** (quando aplicável — especialmente para periódicos brasileiros ou periódicos avaliados no sistema Qualis); **Tempo médio de resposta nos processos editoriais** (tempo médio entre a submissão e a decisão editorial ou revisão, quando disponível); **Motivos da inclusão** e vantagens de submeter o artigo a esse periódico.

Formato da Saída: Organize os resultados em uma tabela clara, em que cada linha representa um periódico e as colunas contêm as informações coletadas sobre ele.

Figura 2.19: *Prompt* para sugestão de locais de publicação.

O *prompt* apresentado na Figura 2.19 foi aplicado com base no presente estudo como artigo-fonte para fins didáticos. Como se pode observar, a lógica de construção é a mesma apresentada ao longo do documento. O resultado obtido é apresentado na na Figura 2.20. Vale ressaltar que, se tratando de um capítulo de livro, a LLM reforça a necessidade de adaptar a escrita e o foco do estudo, e, a partir disso, são sugeridos os periódicos para publicação e apresentadas as devidas justificativas.

Nome do Periódico	Fator de Impacto / Métricas	Estrato Qualis (CAPES 2017-2020)	Tempo Médio (Estimado)	Motivos e Vantagens de Submissão
Communications of the Association for Information Systems (CAIS)	SJR: Q2 (Scopus) H-Index: 68	A2 / A3 (Var. área)	3 a 6 meses	Alinhamento Perfeito: A CAIS é famosa por publicar artigos do tipo "Tutorial" e discussões sobre o campo de SI. Aceita bem papers sobre como fazer pesquisa. Vantagens: Visibilidade global na comunidade de SI; Open Access (sem taxas para membros AIS).
Revista de Administração Contemporânea (RAC)	JCR: 0.8 (JCI) Scielo Index	A2	3 a 5 meses (Fast track disp.)	Seção Específica: Possui a seção de "Artigos Tecnológicos", ideal para propor ferramentas (Research Rabbit, NotebookLM) aplicadas à pesquisa. Vantagens: Uma das melhores do Brasil em Administração/SI; foco em impacto prático; Open Access.
iSys - Revista Brasileira de Sistemas de Informação	H-Index: 12 (Google Scholar)	A3	4 a 8 meses	Público-Alvo: É o periódico oficial da Sociedade Brasileira de Computação (SBC) para SI. Vantagens: O público é exatamente a comunidade de SI citada no seu abstract. Ótima para artigos em português ou inglês com foco na comunidade nacional.
Journal of Information Systems Education (JISE)	SJR: Q2/Q3	A3 / B1	3 a 6 meses	Foco Educacional: Se o seu artigo foca em ensinar pesquisadores/alunos a usar IA (abordagem pedagógica), esta é a melhor revista mundial da área. Vantagens: Especializada em "best practices" e currículo de SI.

Figura 2.20: Resultado requisição para indicação de periódicos

Convém pontuar que o resultado apresentado na Figura 2.20 é um dos mais passíveis de alucinação, principalmente em relação às métricas de impacto, Qualis e ao tempo médio para a primeira decisão. Sendo imprescindível a verificação dos fatos. De toda forma, os resultados obtidos mostraram-se interessantes, considerando a natureza do presente manuscrito. Recomenda-se que os participantes testem utilizando suas próprias produções. Essa e outras recomendações são abordadas na seção a seguir sobre aspectos didáticos do curso.

2.6. Aspectos didáticos

A proposta metodológica do minicurso enfatiza atividades nas quais os participantes utilizam diretamente ferramentas de IA ao longo das diferentes etapas do processo de pesquisa. O curso está sendo ofertado pela primeira vez no formato presencial, com carga horária total de quatro horas, dividida em dois blocos de duas horas cada. A metodologia adotada combina exposições teóricas de curta duração com demonstrações e exercícios práticos orientados. Esse formato visa permitir que os participantes acompanhem e executem as tarefas propostas de forma guiada. As atividades incluem o uso de ferramentas como ResearchRabbit para a expansão da literatura, NotebookLM para a análise e a sumarização de documentos científicos e diferentes LLMs, como GPT, Gemini e DeepSeek, para demonstrar atividades relacionadas à engenharia de *prompts* e à replicação de procedimentos de pesquisa científica. Todo o material produzido encontra-se disponível publicamente no GitHub ¹⁰.

Os autores elaboraram, ministraram e ajustaram a estrutura e o material deste minicurso com base em ofertas anteriores realizadas no formato remoto. Inicialmente, o minicurso foi ministrado a integrantes do Laboratório de Computação Aplicada (LACA) da Universidade Federal do Oeste do Pará (Ufopa), com duração de quatro horas, com o objetivo de avaliar o material didático e o tempo necessário para a realização das atividades. Em seguida, foi ofertado a um grupo de pesquisa composto por discentes de graduação e pós-graduação da Universidade Estadual do Maranhão (UEMA) e da Ufopa, mantendo a mesma carga horária. Posteriormente, o curso foi ministrado como extensão, também em formato remoto, no Instituto de Ciências Matemáticas e de Computação (ICMC) da Universidade de São Paulo, por meio do Laboratório de Inteligência Computacional (LABIC), com duração de seis horas, 150 vagas ofertadas e ampla procura por parte da comunidade acadêmica, com mais de oito mil inscritos.

Visando uma abordagem dialógica e andragógica, encoraja-se a participação ativa dos estudantes por meio do compartilhamento de experiências prévias e de discussões sobre aspectos éticos. A depender do conhecimento prévio dos participantes, um detalhamento técnico das LLMs e das tecnologias adjacentes também se mostra interessante. Dinâmicas com a modificação de *prompts* também são sugeridas, o que pode facilmente torná-lo um curso de 20 a 40 horas, a ser ministrado como disciplina optativa, permitindo que os estudantes apliquem o ferramental aos seus projetos de pesquisa e construam, colaborativamente, *guidelines* e materiais complementares, incorporando novas ferramentas. Vislumbra-se também que avaliações baseadas em relatórios, apresentações e relatos de experiências têm o potencial de enriquecer o processo de ensino-aprendizagem.

¹⁰https://github.com/fabiolobato/cursoiapesquisa_sbsi

2.7. Aspectos Éticos

Além das atividades técnicas, o minicurso incorpora discussões sobre aspectos éticos, de transparência e de responsabilidade no uso de ferramentas de IA generativa no contexto da pesquisa científica. Essa discussão é motivada pelas limitações inerentes às tecnologias baseadas em LLMs, bem como pelos impactos metodológicos decorrentes de seu uso no fazer científico. As ferramentas apresentadas ao longo do minicurso estão sujeitas a limitações técnicas que incluem, entre outras, a possibilidade de geração de respostas imprecisas ou inconsistentes, fenômeno denominado alucinação [Maleki et al. 2024]. Além disso, as tecnologias baseadas em LLMs são constantemente atualizadas, seja por meio de ajustes nos modelos, seja pela alteração de parâmetros, de bases de treinamento ou de políticas de acesso. Isso pode resultar em variações nos resultados obtidos ao longo do tempo, mesmo quando utilizados *prompts* idênticos.

Outro aspecto importante diz respeito aos vieses inerentes aos LLMs, que podem decorrer tanto dos dados de treinamento quanto das escolhas de projeto e das estratégias de ajuste dos modelos. Esses vieses não são apenas dependentes do conteúdo analisado, mas também do próprio sistema, resultando em discrepâncias entre as respostas geradas pelos modelos [Lin et al. 2025]. No contexto da pesquisa científica, esses fatores reforçam a necessidade de uso crítico das ferramentas e de validação humana dos resultados.

Diante deste cenário, as pessoas autoras preconizam a adoção de boas práticas, a fim de mitigar essas limitações e aumentar a transparência e a replicabilidade das pesquisas científicas que utilizam ferramentas de IA generativa. Entre as práticas recomendadas, destacam-se o armazenamento e a documentação dos *prompts* utilizados e das respostas obtidas. Sugere-se incluir metadados sobre a ferramenta, como a versão do modelo e a data e a hora da execução, para facilitar o versionamento. A apresentação pública do que foi gerado por IA, descrevendo os artefatos e como foi realizada a curadoria do material, é uma ação recomendada, pois permite a troca de experiências dentro do grupo de pesquisa e também funciona como um *guardrail* para eventuais falhas não percebidas. Além disso, reforça-se a importância de explicitar, em relatórios e publicações científicas, o papel desempenhado pelas ferramentas de IA no processo de pesquisa, contribuindo para maior transparência e rigor metodológico (ver a subseção “Sobre o uso de IA generativa” após as Considerações Finais).

2.8. Considerações Finais

Neste minicurso foram apresentados conceitos, ferramentas e práticas relacionados ao uso de Inteligência Artificial como suporte ao processo de pesquisa científica, com ênfase em modelos de linguagem de grande escala e em ferramentas atualmente disponíveis para esse processo. Buscou-se integrar fundamentos conceituais às atividades práticas, permitindo que os participantes compreendessem o potencial dessas tecnologias em diferentes etapas do ciclo de pesquisa. Foi proposta uma taxonomia das tarefas de pesquisa que utiliza a IA como co-cientista. As atividades foram dispostas de forma sequencial, considerando a evolução atual da IA agêntica e a aderência ao método científico, e encontram-se agrupadas em quatro fases: revisão da literatura, desenho de pesquisa, condução da pesquisa e publicação. Em cada uma dessas fases, buscou-se apresentar as ferramentas mais adequadas às atividades correspondentes do processo de pesquisa, de modo a destacar os

critérios de uso e de adequação dessas tecnologias em diferentes contextos.

Cabe destacar que o uso de ferramentas envolve desafios relevantes, como as limitações inerentes aos modelos, a variabilidade dos resultados, as atualizações frequentes e os vieses associados às tecnologias baseadas em LLMs. Esses fatores impactam diretamente a confiabilidade, a transparência e a replicabilidade dos processos de pesquisa, exigindo uma postura crítica por parte do pesquisador. Dessa forma, práticas como a validação dos resultados, o registro e o arquivo dos *prompts* e das respostas obtidas e a delimitação explícita do papel das ferramentas de IA nos estudos tornam-se essenciais. O minicurso foi concebido para oferecer um percurso formativo que pode ser incorporado a disciplinas como metodologia científica, ética em pesquisa e produção científica, bem como a componentes curriculares correlatos, tanto no nível de graduação quanto de pós-graduação. Uma versão estendida também é possível, por exemplo, para ser ofertada como disciplina optativa.

Por fim, seguindo os princípios de ciência aberta, todos os *prompts* e os materiais suplementares estão publicamente disponíveis no GitHub: https://github.com/abiolobato/cursoiapesquisa_sbsi, sob licença CC BY-NC 4.0.

Sobre o uso de IA generativa

A ferramenta Gemini 3 Pro foi utilizada para a elaboração da proposta do minicurso nas etapas de ideação, de melhoria do conteúdo e de engenharia de *prompts*. O Grammarly (V. 1.151.1) foi utilizado para a revisão gramatical, já que atualmente também oferece cobertura para o português. Os demais usos foram descritos ao longo do manuscrito. Os textos aqui dispostos são de responsabilidade das pessoas autoras.

Agradecimentos

Este trabalho foi apoiado pelo Conselho Nacional de Desenvolvimento Científico e Tecnológico (CNPq), por meio das bolsas DT-303031/2023-9, PDS-101057/2024-5, PQ-2023/10100-4, bem como pelos auxílios #309575/2021-4 e #307184/2025-0. Parte do trabalho foi realizado no Centro de Inteligência Artificial da Universidade de São Paulo (C4AI – <http://c4ai.inova.usp.br/>), com apoio da Fundação de Amparo à Pesquisa do Estado de São Paulo (FAPESP, processo #2019/07665-4) e da IBM Corporation. O projeto também contou com apoio do Ministério da Ciência, Tecnologia e Inovação, com recursos da Lei n. 8.248, de 23 de outubro de 1991, no âmbito do PPI-SOFTEX, coordenado pela Softex e publicado como Residência em TIC 13, DOU 95, processo 01245.010222/2022-44. Este trabalho recebeu apoio da Coordenação de Aperfeiçoamento de Pessoal de Nível Superior – Brasil (CAPES) – Código de Financiamento 001. Por fim, também recebeu financiamento da Financiadora de Estudos e Projetos (FINEP) – Pró-Amazônia, Referência 2373/24 - CTCCA-II.

Agradecemos às pessoas organizadoras da Trilha de Minicursos do SBSI, Jonice Oliveira e Davi Viana, por todo o apoio ao longo do processo, e às pessoas revisoras pelas valiosas contribuições ao estudo desde a fase de proposta. Agradecemos também a todas as pessoas participantes das versões do minicurso, que serviu de base para a publicação deste trabalho, e que contribuíram com perguntas instigantes, sugestões de melhoria e *feedback*, que muito nos ajudaram na construção deste material.

Referências

- Annepaka, Y. and Pakray, P. (2025). Large language models: a survey of their development, capabilities, and applications. *Knowledge and Information Systems*, 67(3):2967–3022.
- Bahdanau, D., Cho, K., and Bengio, Y. (2014). Neural machine translation by jointly learning to align and translate. *arXiv preprint arXiv:1409.0473*.
- Brown, T., Mann, B., Ryder, N., Subbiah, M., Kaplan, J. D., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., Askell, A., et al. (2020). Language models are few-shot learners. *Advances in neural information processing systems*, 33:1877–1901.
- Chen, Q., Yang, M., Qin, L., Liu, J., Yan, Z., Guan, J., Peng, D., Ji, Y., Li, H., Hu, M., Zhang, Y., Liang, Y., Zhou, Y., Wang, J., Chen, Z., and Che, W. (2025). Ai4research: A survey of artificial intelligence for scientific research.
- DeepMind, G. (2025). Gemini 3 pro. <https://ai.google.dev/gemini-api/docs/models?hl=pt-br>. Multimodal Large Language Model.
- Deng, Z., Ma, W., Han, Q.-L., Zhou, W., Zhu, X., Wen, S., and Xiang, Y. (2025). Exploring deepseek: A survey on advances, applications, challenges and future directions. *IEEE/CAA Journal of Automatica Sinica*, 12(5):872–893.
- Devlin, J., Chang, M.-W., Lee, K., and Toutanova, K. (2019). BERT: Pre-training of deep bidirectional transformers for language understanding. In Burstein, J., Doran, C., and Solorio, T., editors, *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 4171–4186, Minneapolis, Minnesota. Association for Computational Linguistics.
- Gottweis, J., Weng, W.-H., Daryin, A., Tu, T., Palepu, A., Sirkovic, P., Myaskovsky, A., Weissenberger, F., Rong, K., Tanno, R., et al. (2025a). Towards an ai co-scientist. *arXiv preprint arXiv:2502.18864*.
- Gottweis, J., Weng, W.-H., Daryin, A., Tu, T., Palepu, A., Sirkovic, P., Myaskovsky, A., Weissenberger, F., Rong, K., Tanno, R., Saab, K., Popovici, D., Blum, J., Zhang, F., Chou, K., Hassidim, A., Gokturk, B., Vahdat, A., Kohli, P., Matias, Y., Carroll, A., Kulkarni, K., Tomasev, N., Guan, Y., Dhillon, V., Vaishnav, E. D., Lee, B., Costa, T. R. D., Penadés, J. R., Peltz, G., Xu, Y., Pawlosky, A., Karthikesalingam, A., and Natarajan, V. (2025b). Towards an ai co-scientist.
- Jones, K. S. (1994). Natural language processing: a historical review. *Current issues in computational linguistics: in honour of Don Walker*, pages 3–16.
- Joseph, S. R., Hlomani, H., Letsholo, K., Kaniwa, F., and Sedimo, K. (2016). Natural language processing: A review. *International Journal of Research in Engineering and Applied Sciences*, 6(3):207–210.

- Liddy, E. D. (2001). Natural language processing. In *Encyclopedia of Library and Information Science*. Marcel Dekker, Inc., New York, 2nd edition.
- Lin, L., Wang, L., Guo, J., and Wong, K.-F. (2025). Investigating bias in llm-based bias detection: Disparities between llms and human perception. In *Proceedings of the 31st International Conference on Computational Linguistics*, pages 10634–10649.
- Maleki, N., Padmanabhan, B., and Dutta, K. (2024). Ai hallucinations: a misnomer worth clarifying. In *2024 IEEE conference on artificial intelligence (CAI)*, pages 133–138. IEEE.
- Mienye, I. D., Swart, T. G., and Obaido, G. (2024). Recurrent neural networks: A comprehensive review of architectures, variants, and applications. *Information*, 15(9):517.
- Mishra, T., Sutanto, E., Rossanti, R., Pant, N., Ashraf, A., Raut, A., Uwabareze, G., Oluwatomiwa, A., and Zeeshan, B. (2024). Use of large language models as artificial intelligence tools in academic research and publishing among global clinical researchers. *Scientific reports*, 14.
- OpenAI (2025). Gpt-5. <https://platform.openai.com/docs/models>. Large Language Model.
- Peters, M. E., Neumann, M., Iyyer, M., Gardner, M., Clark, C., Lee, K., and Zettlemoyer, L. (2018). Deep contextualized word representations. In Walker, M., Ji, H., and Stent, A., editors, *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long Papers)*, pages 2227–2237, New Orleans, Louisiana. Association for Computational Linguistics.
- Peykani, P., Ramezanlou, F., Tanasescu, C., and Ghanidel, S. (2025). Large language models: A structured taxonomy and review of challenges, limitations, solutions, and future directions. *Applied Sciences*, 15(14):8103.
- Radford, A., Narasimhan, K., Salimans, T., Sutskever, I., et al. (2018). Improving language understanding by generative pre-training.
- Radford, A., Wu, J., Child, R., Luan, D., Amodei, D., Sutskever, I., et al. (2019). Language models are unsupervised multitask learners. *OpenAI blog*, 1(8):9.
- Si, C., Yang, D., and Hashimoto, T. (2024). Can llms generate novel research ideas? a large-scale human study with 100+ nlp researchers. *13th International Conference on Learning Representations, ICLR 2025*, pages 56923–57012.
- Toosi, A., Bottino, A. G., Saboury, B., Siegel, E., and Rahmim, A. (2021). A brief history of ai: how to prevent another winter (a critical review). *PET clinics*, 16(4):449–469.
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł., and Polosukhin, I. (2017). Attention is all you need. *Advances in neural information processing systems*, 30.

Wang, Z., Chu, Z., Doan, T. V., Ni, S., Yang, M., and Zhang, W. (2025). History, development, and principles of large language models: an introductory survey. *AI and Ethics*, 5(3):1955–1971.

Wu, T., He, S., Liu, J., Sun, S., Liu, K., Han, Q.-L., and Tang, Y. (2023). A brief overview of chatgpt: The history, status quo and potential future development. *IEEE/CAA Journal of Automatica Sinica*, 10(5):1122–1136.