



SBSI 2026

22ª Edição - Vitória - Espírito Santo

XXII Simpósio Brasileiro de Sistemas de Informação

Sistemas de Informação Inteligentes:
Inovações, Aplicações e Ética na Inteligência Artificial

MINICURSOS EM SI

TÓPICOS ESPECIAIS EM SISTEMAS DE INFORMAÇÃO

Organizadores:

Jonice Oliveira (UFRJ)

Davi Viana (UFMA)



Organizadores da Trilha

Jonice Oliveira

Davi Viana

TÓPICOS ESPECIAIS EM SISTEMAS DE INFORMAÇÃO
Trilha de Minicursos em Sistemas de Informação
SBSI 2026

<https://sbsi.sbc.org.br/2026/>

<https://books-sol.sbc.org.br/index.php/sbc/catalog/category/si>

Sociedade Brasileira de Computação
Porto Alegre
2026



TÓPICOS ESPECIAIS EM SISTEMAS DE INFORMAÇÃO

Trilha de Minicursos em Sistemas de Informação

SBSI 2026

Sociedade Brasileira de Computação -SBC
CNPJ: 29.532.264/0001-78

Coordenação de Publicações

Paulo Malcher (UFRA)
Luiz Paulo Carvalho (UFRRJ)
Rodrigo Zacarias (UFF)
Washington Cunha (UNICAMP)

Coordenação da Trilha de Minicursos em Sistemas de Informação

Jonice Oliveira (UFRJ)
Davi Viana (UFMA)

Coordenação Geral

Karin Satie Komati (IFES)
Vítor Estêvão Silva Souza (UFES)

Editora

Sociedade Brasileira de Computação - SBC

Realização



Organização



Patrocínio



HOTEL SENAC
Ilha do Boi



Apoio



Editores

Paulo Malcher (UFRA)
Luiz Paulo Carvalho (UFRRJ)
Rodrigo Zacarias (UFF)
Washington Cunha (UNICAMP)
Jonice Oliveira (UFRJ)
Davi Viana (UFMA)
Karin Satie Komati (IFES)
Vitor Estêvão Silva Souza (UFES)

Comitê Técnico

Coordenação dos Anais

Paulo Malcher (UFRA)
Luiz Paulo Carvalho (UFRRJ)
Rodrigo Zacarias (UFF)
Washington Cunha (UNICAMP)

Coordenação da Trilha de Minicursos em Sistemas de Informação

Jonice Oliveira (UFRJ)
Davi Viana (UFMA)

Coordenação Geral

Karin Satie Komati (IFES)
Vitor Estêvão Silva Souza (UFES)

Membros da Comissão Especial de Sistemas de Informação da SBC (CE-SI)

Coordenação da CE-SI (2025-2026)

Célia Ghedini Ralha (UnB e UFBA) - Coordenadora CE-SI

Claudia Cappelli Alo (UERJ) - Coordenadora Adjunta

Flávia Maria Santoro (UERJ e Inteli) - Coordenadora Adjunta

Comitê Gestor da CE-SI (2025-2026)

Célia Ghedini Ralha (UnB and UFBA)

Claudia Cappelli Alo (UERJ)

Daniela Barreiro Claro (UFBA)

Flávia Maria Santoro (UERJ e Inteli)

Jonice de Oliveira Sampaio (UFRJ)

José Maria Nazar David (UFJF)

Karin Satie Komati (IFES)

Paulo Robson Campelo Malcher (UFRA)

Rodrigo Pereira dos Santos (UNIRIO)

Sean Wolfgang Matsui Siqueira (UNIRIO)

Valdemar Vicente Graciano Neto (UFG)

Comitê de Programa - Trilha de Minicursos em Sistemas de Informação do SBSI 2026

Alessandro Cerqueira	UniLaSalle-RJ / FAETERJ
Alex Borges	UFJF
André Freire	UFLA
André Martinotto	UCS
Arlino Magalhães	UFPI
Awdren Fontão	UFMS
Carlos Eduardo Pires	UFCG
Carlos Eduardo de Barros Paes	PUC-SP
Cristiano da Silveira Colombo	IFES
Daniel Notari	UCS
Emanuel Coutinho	UFC
Fatima Nunes	EACH-USP
Fernanda Maria Santos	UFU
Flavia Maria Santoro	UERJ
Flávio Soares Corrêa da Silva	USP
Francisco Henrique Ferreira	UFJF
Glauco Carneiro	UFS
Heitor Costa	UFLA
Jonice Oliveira	UFRJ
Jorge Barbosa	UNISINOS
José Maria David	UFJF
Juliano Lopes de Oliveira	UFG
Leila Bergamasco	FEI
Maria Istela Cagnin	UFMS
Paulo Malcher	UFRA
Paulo Sérgio Santos	UNIRIO
Rafael Testa	UNIVESP
Ricardo Choren	IME
Silas Lima Filho	UFRJ
Victor Stroele	UFJF



Esta obra está sob a licença Creative Commons Atribuição 4.0 (CC-BY). Você pode redistribuir este livro em qualquer suporte ou formato e copiar, remixar, transformar e criar a partir do conteúdo deste livro para qualquer fim, desde que cite a fonte.

Dados Internacionais de Catalogação na Publicação (CIP)

S612 Simpósio Brasileiro de Sistemas de Informação (22. : 25 – 28 maio 2026 : Vitória)

Minicursos do SBSI 2026 [recurso eletrônico] / organização: Jonice Oliveira e Davi Viana. – Dados eletrônicos. – Porto Alegre: Sociedade Brasileira de Computação, 2026.

84 p. : il. : PDF ; 1 MB

Modo de acesso: World Wide Web.

ISBN 978-85-7669-665-0 (e-book)

1. Computação – Brasil – Evento. 2. Sistemas de informação. 3. Inteligência artificial. I. Oliveira, Jonice. II. Viana, Davi. III. Sociedade Brasileira de Computação. VI. Título.

CDU 004(063)

Ficha catalográfica elaborada por Annie Casali – CRB-10/2339

Biblioteca Digital da SBC – SBC OpenLib



Sociedade Brasileira de Computação

Av. Bento Gonçalves, 9500

Setor 4 | Prédio 43.412 | Sala 219 | Bairro

Agronomia Caixa Postal 15012 | CEP 91501-970

Porto Alegre - RS

(51) 99252-6018

sbc@sbc.org.br

Prefácio:

Comissão Especial de Sistemas de Informação (CESI) da Sociedade Brasileira de Computação (SBC)

Prezados participantes do XXII Simpósio Brasileiro de Sistemas de Informação (SBSI 2026),

Em nome do Comitê Gestor da Comissão Especial de Sistemas de Informação da Sociedade Brasileira de Computação (CESI/SBC), gostaríamos de dar as boas-vindas a mais uma edição do SBSI. A CESI congrega pesquisadores e profissionais de Sistemas de Informação (SI) e realiza a gestão desta área no Brasil no contexto da SBC. O seu trabalho consiste em articular políticas para divulgação, fortalecimento, consolidação e melhoria da qualidade da educação, pesquisa, inovação e atuação em SI no Brasil, promovendo e propondo importantes atividades para o desenvolvimento e consolidação desta área em nosso país. O SBSI é uma das principais ações da CESI. Por esta razão, ficamos felizes que autores submetam seus trabalhos, que avaliadores se esforcem para dar o melhor feedback possível em suas revisões e que coordenadores consigam realizar momentos memoráveis para que, assim, a comunidade se fortaleça, atraia mais participantes e se consolide no país. Neste sentido, a CESI convoca a comunidade de SI a ser cada vez mais ativa e os participantes do SBSI a se envolverem e trazerem novos pesquisadores para o evento, fazendo a comunidade crescer e solidificar as suas pesquisas.

O SBSI é o fórum mais importante do Brasil na área de SI e reúne anualmente pesquisadores de todo o país. Este ano, o SBSI aconteceu em Vitória, Espírito Santo, sob a coordenação geral dos professores Karin Satie Komati (IFES) e Vítor Estêvão Silva Souza (UFES). Após vários meses de trabalho e dedicação dos professores e de sua equipe de organização local, disponibilizamos à comunidade de SI uma programação rica, diversificada, abrangente, e sólida, de forma a comunicar os resultados mais recentes da pesquisa na área de SI no Brasil. O evento tem crescido consideravelmente nos últimos anos. Tal feito é reflexo da consolidação da área, que somente tem sido possível devido ao apoio crescente de toda a comunidade de SI. O SBSI 2026 é composto por várias atividades em paralelo, que incluem as sessões técnicas, palestras, painéis, master classes e cinco trilhas: Trilha de Pesquisa em SI; Trilha de Indústria e Inovação em SI; Trilha de Temas, Ideias e Resultados Emergentes em SI; Concurso de Teses, Dissertações e Trabalhos de Graduação em SI; e Trilha de Minicursos em SI. Além disso, o Seminário de Grandes Desafios de Pesquisa em

Sistemas de Informação no Brasil para os próximos dez anos (GranDSI-BR 2026-2036) inclui apresentações de projetos. Todos os trabalhos aceitos e apresentados nas demais trilhas além da Trilha de Pesquisa em SI são publicados nos Anais Estendidos do SBSI, também disponibilizados abertamente para a comunidade, e o material didático na forma de capítulos oriundos dos minicursos são publicados como um e-book.

O Comitê de Programa da Trilha de Minicursos em SI, coordenado pelos professores Jonice Oliveira (UFRJ) e Davi Viana (UFMA), trabalhou arduamente e de forma exímia para que os trabalhos mais maduros e de maior qualidade pudessem ser trazidos à comunidade brasileira de SI. Tal trabalho de seleção rigoroso foi fruto de uma confiança plena na qualidade dos trabalhos veiculados nesse simpósio em 2026. O impacto do trabalho realizado pelo comitê de programa e pelos autores dos capítulos de minicursos, potencializando a abrangência e o alcance da produção científica do SBSI, também é observado pelo fato dos anais do SBSI serem indexados e acessíveis mundialmente pela base DBLP. Em tempo, é importante ressaltar que os anais do SBSI, os anais estendidos do SBSI e o livro de minicursos em SI estão sendo indexados e disponibilizados no Portal de Publicações e Conteúdo Digital da SBC (SBC OpenLib - SOL).

Por fim, agradecemos a todos os participantes, colaboradores, coordenadores, patrocinadores, apoiadores e organizadores que, mais uma vez, tornaram possível a realização deste evento presencial.

Desejamos a todos um excelente SBSI 2026!

Um grande abraço,

Célia Ghedini Ralha (UnB and UFBA)
Coordenadora da CESI/SBC - 2025/2026

Flávia Maria Santoro (UERJ and Inteli)
Coordenadora Adjunta da CESI/SBC - 2025/2026

Claudia Cappelli (UERJ)
Coordenadora Adjunta da CESI/SBC - 2025/2026



Célia Ghedini Ralha possui doutorado pela University of Leeds, Inglaterra (1996), mestrado em Engenharia Eletrônica e Computação pelo Instituto Tecnológico de Aeronáutica (ITA) (1990) e pós-doutorado na Brandenburgische Technische Universität Cottbus, Alemanha (2012-2013). Foi professora da Universidade de Brasília (UnB), no Departamento de Ciência da Computação (2002-2017), onde atuou como coordenadora do curso de Bacharelado em Ciência da Computação (2004-2006), chefe de departamento (2006-2010) e diretora de Desenvolvimento Institucional e Inovação no Decanato de Pesquisa e Pós-Graduação (2010-2012). É membro do Programa de Pós-Graduação em Informática da UnB desde 2003, tendo coordenado o

programa entre 2015-2017, classificado com nota 6 pela CAPES (2021-2024). Foi professora visitante na Australian National University (ANU) e no Programa de Pós-Graduação em Ciência da Computação da Universidade Federal da Bahia (UFBA) (2024-2026). Célia Ralha é bolsista de produtividade do Conselho Nacional de Desenvolvimento Científico e Tecnológico (CNPq) desde 2015. Sua experiência está na área de Ciência da Computação, com foco em sistemas inteligentes, especialmente modelos baseados em agentes, sistemas multiagentes, planejamento e simulação baseada em agentes. Ela lidera o grupo de pesquisa InfoKnow (Sistemas Computacionais para Tratamento de Informação e Conhecimento), credenciado pelo CNPq desde 2010. Célia Ralha é membro sênior da Sociedade Brasileira de Computação (SBC) e atuou como coordenadora-geral da Comissão Especial de Sistemas de Informação (CESI) nos períodos 2013-2015, 2024-2025 e 2025-2026, além de vice-coordenadora em 2010-2011. Desde 2024, é editora-chefe da Revista Eletrônica de Iniciação Científica em Computação (REIC), da SBC, e desde 2022 atua como editora associada do periódico Knowledge and Information Systems (KAIS), da Springer.



Flavia Maria Santoro é Diretora Acadêmica do Instituto de Tecnologia e Liderança (Inteli) e professora da Universidade do Estado do Rio de Janeiro (UERJ). Ela obteve seu doutorado em Engenharia de Sistemas e Computação pela Universidade Federal do Rio de Janeiro (UFRJ), além de possuir graduação em Engenharia Eletrônica pela Escola Politécnica da UFRJ e mestrado em Filosofia Contemporânea pela Pontifícia Universidade Católica do Rio de Janeiro (PUC-Rio). Desde 2009, é bolsista de produtividade do Conselho Nacional de Desenvolvimento Científico e Tecnológico (CNPq). Sua trajetória acadêmica também inclui períodos sabáticos na Université Pierre et Marie Curie - Paris VI, na França, e na Queensland University of Technology, na Austrália. Com mais de duas décadas de experiência como educadora e pesquisadora na área de

Sistemas de Informação, seus trabalhos concentram-se principalmente em Gestão de Processos de Negócio (BPM), Processos Intensivos em Conhecimento, Gestão do Conhecimento, Trabalho Cooperativo e Aprendizagem Suportados por Computador (CSCW/CSCL), e Ética em BPM. Além da atuação acadêmica, tem contribuído como consultora em diversos projetos relacionados a BPM e desenvolvimento de software

para diferentes empresas. Atualmente, é Diretora de Inovação da Sociedade Brasileira de Computação (SBC) no período de 2025 a 2027.



Claudia Cappelli é professora da Universidade do Estado do Rio de Janeiro (UERJ). Também atua como professora no Mestrado Profissional em Telemedicina e Telessaúde da UERJ e colabora com o Programa de Pós-Graduação em Sistemas de Informação da Universidade Federal do Rio de Janeiro (UFRJ). Possui doutorado em Ciências - Informática pela Pontifícia Universidade Católica do Rio de Janeiro (PUC-Rio), graduação em Sistemas de Informação pela UERJ e mestrado também em Sistemas de Informação pela UFRJ. Realizou pós-doutorado no Programa de Pós-Graduação em

Informática da UNIRIO (2010) e na UFRJ (2020). É fundadora do Instituto Nacional de Ciência e Tecnologia em Democracia Digital (INCT-DD) e pesquisadora em Letramento Digital no IBCIH (Instituto Brasileiro de Cidades Inteligentes e Humanas). Atua como gestora do conhecimento no Linguagem Simples Lab e na Rede Brasileira de Linguagem Simples. É membro do comitê gestor do programa Meninas Digitais da Sociedade Brasileira de Computação (SBC), representante brasileira na organização Clarity e membro da PLAIN (Plain Language Association International). Também é pesquisadora em inovação no INEI (Instituto Nacional de Empreendedorismo e Inovação). Na indústria, foi gerente da área de Arquitetura Corporativa e Planejamento de Tecnologia no Citibank (1996-2001) e na Telemar (2001-2003). Atuou ainda por 8 anos na Petrobras em um projeto de Gestão de Processos de Negócio (2008-2016) e por 2 anos na Caixa Econômica Federal (CEF) em Ciência de Dados com foco em Linguagem Simples (2018-2020). Exerceu o cargo de Diretora de Relações com a Indústria da SBC (2016-2018) e atua como revisora de diversos periódicos nacionais e internacionais. Sua atuação está na área de Sistemas de Informação, com ênfase em Gestão de Processos de Negócio, Arquitetura Empresarial, Gestão de TI, Transparência, Linguagem Simples e Governo Digital.

Prefácio:

Trilha de Minicursos em Sistemas de Informação do SBSI 2026

Este livro reúne os minicursos do XXII Simpósio Brasileiro de Sistemas de Informação (SBSI 2026). O evento contou com a organização compartilhada entre o Instituto Federal do Espírito Santo (IFES) e a Universidade Federal do Espírito Santo (UFES), no período de 25 a 28 de maio de 2026. Participam do SBSI, fórum nacional de debates da área de Sistemas de Informação (SI), estudantes e pesquisadores com apresentação de trabalhos científicos e discussão de temas contemporâneos relacionados à área. Esta edição foi realizada presencialmente. O SBSI 2026 estabeleceu como temática principal **"Sistemas de Informação Inteligentes: Inovações, Aplicações e Ética na Inteligência Artificial"**. A escolha desta temática reflete a urgência e a importância de se discutir como a inteligência artificial impacta inovações e ética em sistemas de informação.

Neste ano, foram submetidas 9 (nove) propostas de minicursos válidas e 2 (duas) foram selecionadas, tal qual decidido em reunião da Comissão Especial de Sistemas de Informação e definido na chamada de minicursos. Tais propostas foram avaliadas por, no mínimo, três professores doutores que fizeram parte do comitê científico da trilha de Minicursos do SBSI 2026.

Os dois minicursos contidos neste livro abordam tópicos de interesse da comunidade de Sistemas de Informação. O primeiro capítulo, intitulado **"Sistemas de Informação Além do Sociotécnico: Explorando o Entrelaçamento entre IA e Humanos por meio do Design Especulativo"**, discute os limites do paradigma sociotécnico tradicional para compreender Sistemas de Informação em contextos marcados pela integração crescente da Inteligência Artificial às práticas humanas. A partir de perspectivas críticas, como o conceito de ciborgue e as teorias do emaranhamento sociotécnico, os autores propõem uma leitura relacional das interações entre humanos, tecnologias e outros actantes. Como contribuição, apresentam a abordagem *Speculative Entangled Design (SpED)* e o *Sociotechnical Entanglement Framework (SEF)*, oferecendo ferramentas para refletir criticamente sobre o design de tecnologias e seus impactos éticos, sociais e políticos

O segundo capítulo, intitulado **"Ferramentas de Inteligência Artificial para suporte ao processo de pesquisa em Sistemas de Informação"**, apresenta como a ascensão da Inteligência Artificial, especialmente dos Large Language Models (LLMs), está transformando o processo de pesquisa científica e criando um novo ecossistema de ferramentas de apoio ao trabalho acadêmico. Os autores apresentam, de forma prática, diferentes plataformas e técnicas que auxiliam etapas do ciclo de pesquisa, incluindo revisão de literatura, desenho metodológico, análise de artigos e escrita científica. Ao mesmo tempo, o capítulo enfatiza o uso crítico e responsável dessas tecnologias, destacando desafios éticos, de transparência e de reprodutibilidade no uso da IA na produção científica.

Esperamos que este livro auxilie estudantes, pesquisadores e profissionais da área de Sistemas de Informação na construção do conhecimento em temas específicos relacionados ao que foi aqui apresentado. Para alguns, poderá apresentar novos assuntos vinculados à sua área de atuação. Para outros, oferecerá um aprofundamento em temas já conhecidos ou a oportunidade de extrair elementos a serem aplicados em sua pesquisa e/ou prática.

Jonice Oliveira (UFRJ) & Davi Viana (UFMA)
Coordenação da Trilha de Minicursos em Sistemas de Informação



Jonice Oliveira é professora associada do Instituto de Computação da Universidade Federal do Rio de Janeiro (UFRJ) e membro efetivo do Programa de Pós-Graduação em Informática (PPGI/UFRJ). Criou e coordena o Laboratório CORES (Laboratório de Computação Social e Análise de Redes Sociais), que conduz pesquisas multidisciplinares para o entendimento, simulação e fomento às interações sociais. Suas principais áreas de pesquisa são Ciência de Dados e Computação Social, com foco em Big Social Data (tratamento, gestão e extração de conhecimento).



Davi Viana é professor adjunto da Universidade Federal do Maranhão (UFMA). É bolsista de Produtividade em Pesquisa do CNPq (Nível C). Doutor e mestre em Informática pela Universidade Federal do Amazonas (UFAM). Graduado em Ciência da Computação pela Universidade Federal do Amazonas (UFAM). Tem interesse nas áreas de Sistemas de Informação, qualidade de software, melhoria de processos de software (SPI), programas com ênfase na adoção de modelos de maturidade, Educação em Engenharia de Software, Engenharia de Software Experimental, Sistemas de Informação, Engenharia de Software para Internet das Coisas e startups de software.

Prefácio:

Coordenação Geral do SBSI 2026

É um prazer recebê-los no 22º Simpósio Brasileiro de Sistemas de Informação (SBSI 2026). Neste ano, o principal fórum de Sistemas de Informação no Brasil acontece em Vitória, Espírito Santo, de 25 a 28 de maio de 2026. Conhecida por seu litoral deslumbrante e por seu crescente ecossistema tecnológico, Vitória oferece o cenário ideal para que nossa comunidade se conecte, compartilhe conhecimento e projete o futuro da nossa área no Hotel Senac Ilha do Boi.

Organizada pelo Instituto Federal do Espírito Santo (IFES), em colaboração com a Universidade Federal do Espírito Santo (UFES) e a Universidade Vila Velha (UVV), sob os auspícios da Comissão Especial de Sistemas de Informação (CESI) da Sociedade Brasileira de Computação (SBC), esta edição dá continuidade à tradição de excelência do SBSI. Reunimos uma combinação vibrante de pesquisadores de destaque, profissionais da indústria, estudantes e praticantes de todo o país e do exterior, promovendo um ambiente verdadeiramente colaborativo.

No centro do SBSI 2026 está o tema “Sistemas de Informação Inteligentes: Inovações, Aplicações e Ética na Inteligência Artificial”. À medida que modelos e ferramentas de IA se tornam cada vez mais ubíquos, cresce também nossa responsabilidade como profissionais de SI. O tema deste ano nos convida a explorar não apenas o poder transformador da IA em setores como saúde, educação, indústria e gestão pública, mas também as dimensões éticas cruciais, como transparência, privacidade, vieses e impacto social, que devem orientar esses avanços tecnológicos.

Ao longo de quatro dias de intensa troca, os participantes vivenciarão uma programação rica, com sessões técnicas, palestras nacionais e internacionais e trilhas especializadas que abrangem todo o espectro da pesquisa e da prática em SI. Temos a honra de receber palestrantes de destaque, como o professor Giancarlo Guizzardi (University of Twente) e a professora Ana Cristina Bicharra Garcia (UNIRIO). Além disso, temos orgulho de destacar iniciativas como o momento SI4All, que reforçam nosso compromisso com a diversidade, a equidade e a inclusão de minorias na área da computação.

Um evento dessa magnitude é resultado do esforço coletivo de uma comunidade dedicada. Expressamos nossa mais profunda gratidão à comissão organizadora, aos coordenadores de trilhas, aos revisores e aos estudantes voluntários, cujo trabalho incansável tornou este simpósio possível. Agradecemos também aos nossos patrocinadores e apoiadores por sua contribuição essencial para o sucesso deste evento.

Convidamos você a participar ativamente das sessões, explorar as belezas de Vitória e, acima de tudo, construir as parcerias que irão moldar a próxima década dos Sistemas de Informação no Brasil.

Bem-vindos ao SBSI 2026!

Cordialmente,

Karin Komati (IFES) e Vítor E. Silva Souza (UFES)
Coordenadores Gerais do SBSI 2026



Karin Satie Komati é doutora em Engenharia Elétrica e mestre em Informática pela Universidade Federal do Espírito Santo (UFES), onde também se graduou em Ciência da Computação e em Engenharia Elétrica. Professora desde 1998, tem sido frequentemente homenageada por turmas concluintes como paraninfa e professora destaque ao longo de sua carreira. Atualmente, é professora do Instituto Federal do Espírito Santo (Ifes) - Campus Serra, atuando no curso de graduação em Sistemas de Informação e no Mestrado em Computação Aplicada. Sua trajetória de liderança no Ifes inclui dois mandatos como Diretora de Pesquisa, Pós-Graduação e Extensão (DPPGE), além de ter exercido as

funções de Coordenadora de Pesquisa e Coordenadora de Programa de Pós-Graduação. Seu impacto institucional é evidenciado por sua liderança na criação dos cursos de mestrado em Engenharia de Controle e Automação e em Computação Aplicada, além de sua contribuição para a proposta bem-sucedida de doutorado em Engenharia Metalúrgica e de Materiais. Sua pesquisa concentra-se em Processamento Digital de Imagens, Visão Computacional, Ciência de Dados e Aprendizado de Máquina, e ela é a pesquisadora responsável pelo grupo de pesquisa Nu[Tec]². Atuante na comunidade científica, foi Vice-Coordenadora do Fórum de Coordenadores de Pós-Graduação da SBC (Sociedade Brasileira de Computação) no biênio 2023-2024 e, desde 2025, ocupa o cargo de Vice-Secretária Regional da SBC no estado do Espírito Santo. Sua produção científica ultrapassa 200 publicações e a orientação de mais de 130 estudantes. Ela liderou ou participou de mais de 70 projetos e recebeu mais de 15 premiações em eventos técnico-científicos. Sua forte articulação com o setor produtivo inclui parcerias com grandes organizações, como ArcelorMittal e Mogai, com apoio de agências de fomento como CAPES, CNPq e FAPES. Além disso, atuou em 13 ocasiões como assessora da SETEC/MEC e atualmente é Editora-Chefe da revista Ifes Ciência.



Vítor E. Silva Souza obteve o título de mestre pela Universidade Federal do Espírito Santo (UFES), no Brasil, e o título de doutor pela Universidade de Trento, na Itália. Vítor é professor associado do Departamento de Informática do Centro Tecnológico da UFES, atuando nos cursos de graduação em Ciência da Computação e Engenharia da Computação, bem como no Programa de Pós-Graduação em Informática (mestrado e doutorado). Na UFES, lidera o Laboratório de Práticas em Engenharia de Software (LabES) e coordena projetos de extensão com a participação de diversos estudantes de graduação. Também é pesquisador sênior do Núcleo de Estudos em Modelagem Conceitual e Ontologias (NEMO), desenvolvendo pesquisas e orientando alunos

nas áreas de Engenharia de Ontologias, Modelagem Conceitual e Engenharia de

Software, com interesse especial em Engenharia de Requisitos, Engenharia Web e publicação de Dados Conectados na chamada Web Semântica. Na Sociedade Brasileira de Computação (SBC), Vítor atuou junto à Comissão de Educação como coordenador do Fórum de Coordenadores de Pós-Graduação entre 2019 e 2021 e, desde 2023, exerce o cargo de Secretário Regional da SBC no estado do Espírito Santo, Brasil.

Sumário

Capítulo 1. Sistemas de Informação Além do Sociotécnico: Explorando o Entrelaçamento entre IA e Humanos com Design Especulativo.....01-32

Marcelo Soares Loutfi, Gabriel de Melo Guedes Souza, Fabiano da Fonseca Ramos, Thaís de Souza Deluca Ferreira, Sean Wolfgang Matsui Siqueira

Capítulo 2. Ferramentas de Inteligência Artificial para suporte ao processo de pesquisa em Sistemas de Informação.....33-63

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

Capítulo

1

Sistemas de Informação Além do Sociotécnico: Explorando o Entrelaçamento entre IA e Humanos com Design Especulativo

*Information Systems Beyond the Sociotechnical:
Exploring the Entanglement Between AI and
Humans through Speculative Design*

Marcelo Soares Loutfi, Gabriel de Melo Guedes Souza, Fabiano da Fonseca Ramos, Thaís de Souza Deluca Ferreira, Sean Wolfgang Matsui Siqueira

Abstract

This chapter presents a critical reflection on the limits of the traditional understanding of Information Systems grounded in the classical sociotechnical paradigm, arguing that this perspective becomes insufficient in the face of the increasingly deep integration of Artificial Intelligence into human practices. The text shows how the separation between technology and society tends to obscure the relational, hybrid, and politically situated nature of contemporary technologies. In this context, Donna Haraway's figure of the cyborg is mobilized as a central conceptual lens to challenge modern dichotomies and to provide a critical basis for understanding humans, technologies, and other actors as co-constitutive entities within sociotechnical contexts. Speculative Design, in turn, is mobilized as an approach capable of critically exploring possible futures, allowing the ethical, social, and political implications of technologies to be problematized before they stabilize in the social world. Based on this theoretical framework, the Speculative Entangled Design (SpED) approach is presented, grounded in Entanglement Theories and making explicit its ethical, epistemological, and ontological commitments. SpED is operationalized through the Sociotechnical Entanglement Framework (SEF) and its set of tools, highlighting how these tools support the formulation of theoretically informed and ethically situated technological requirements.

Resumo

Este capítulo apresenta uma reflexão crítica sobre os limites da compreensão tradicional dos Sistemas de Informação, fundamentada no paradigma sociotécnico clássico, argumentando que essa perspectiva se torna insuficiente diante da integração da Inteligência Artificial às práticas humanas. O texto demonstra como a separação entre tecnologia e sociedade tende a obscurecer a natureza relacional, híbrida e politicamente situada das tecnologias contemporâneas. Nesse contexto, a figura do ciborgue, de Donna Haraway, é mobilizada como uma lente conceitual central para desafiar as dicotomias modernas e fornecer uma base crítica para a compreensão de humanos, tecnologias e outros atores como entidades coconstitutivas em contextos sociotécnicos. O Design Especulativo, por sua vez, é mobilizado como uma abordagem capaz de explorar criticamente futuros possíveis, permitindo que as implicações éticas, sociais e políticas das tecnologias sejam problematizadas antes de se estabilizarem no mundo social. Com base nesse arcabouço teórico, apresenta-se a abordagem do Speculative Entangled Design (SpED), fundamentada nas Teorias do Emaranhado, explicitando seus compromissos éticos, epistemológicos e ontológicos. A SpED é operacionalizada por meio da Sociotechnical Entanglement Framework (SEF) e de seu conjunto de ferramentas, destacando como essas ferramentas apoiam a formulação de requisitos tecnológicos teoricamente informados e eticamente situados.

1.1. Introdução

Já parou para refletir sobre como a Inteligência Artificial (IA) está, silenciosamente, redefinindo nossas decisões, relações e modos de existir? Este minicurso parte justamente dessa inquietação para convidar você a sair do lugar comum e explorar, de forma crítica e criativa, os futuros possíveis dos Sistemas de Informação (SI). Em vez de tratar a tecnologia como neutra ou inevitável, a proposta é provocar questionamentos, ampliar repertórios e experimentar novas formas de pensar o *design* de tecnologia em um mundo profundamente atravessado pela IA.

Durante muito tempo, aprendemos a ver o mundo dividido em duas partes bem separadas. De um lado, a perspectiva dos fatos objetivos, daquilo que pode ser medido, calculado e observado. Do outro, a perspectiva da sociedade, dos valores, das interpretações e das decisões humanas. Essa forma de organizar o pensamento faz com que o mesmo fenômeno seja frequentemente analisado de maneiras distintas, como se as dimensões objetiva e social existissem independentemente.

Nessa forma de separar a realidade, a mudança climática, por exemplo, pode ser analisada por meio de dados mensuráveis, como o aumento da temperatura média global, a concentração de CO₂ na atmosfera, a elevação do nível do mar e modelos matemáticos de previsão climática. Por outro lado, a mudança climática também pode ser vista como um problema político, econômico e moral, envolvendo disputas entre países, impactos desiguais sobre populações vulneráveis, negação científica e o conflito entre crescimento econômico e sustentabilidade.

Por sua vez, na Computação, um algoritmo costuma ser estudado como um artefato técnico, fundamentado em estatísticas, aprendizado de máquina, métricas de pre-

cisão, desempenho e eficiência computacional. Já nos estudos de Ciência, Tecnologia e Sociedade (CTS), o algoritmo é analisado em termos de seus efeitos sociais, como a criação de bolhas informacionais, o reforço de vieses, a manipulação da atenção e os impactos culturais e políticos.

Essa forma de pensar, conhecida como paradigma da bifurcação, tem origem no dualismo cartesiano, segundo o qual a realidade é composta por duas substâncias independentes. A *res extensa* corresponde ao mundo físico e material e a *res cogitans* compreende o pensamento, a consciência e a subjetividade. A ciência moderna herdou essa cisão ao atribuir centralidade ao sujeito cognoscente e ao tratar a natureza como um domínio exterior, passivo e disponível à observação, ao controle e à manipulação.

No século passado, Whitehead (1920) já adotava uma postura crítica em relação ao paradigma da bifurcação. Para ele, essa divisão é artificial; os fenômenos que compõem o mundo não podem ser separados em realidades matematicamente puras e experiências subjetivas confinadas ao domínio privado. Ambos constituem aspectos inseparáveis de processos relacionais que dão forma à natureza dos fenômenos. Bruno Latour aprofunda essa crítica ao demonstrar que a modernidade nunca conseguiu sustentar a suposta pureza da divisão entre natureza e sociedade. Para ele, a ideia de purificação dos campos é uma falácia, pois esses domínios são produzidos simultaneamente por práticas que hibridizam o mundo [Latour 2012].

O paradigma da bifurcação aparece, por exemplo, na própria organização das disciplinas acadêmicas: conteúdos como algoritmos e programação tendem a ser classificados como “técnicos” ou “exatos”, enquanto áreas como Informática e Sociedade ou ética são tratadas como “campos externos”, “sociais” ou “complementares”; como se não integrassem o mesmo ecossistema sociotécnico. Esse descompasso entre disciplinas técnicas e sociotécnicas reforça a bifurcação moderna ao isolar técnica e sociedade em domínios distintos, quando, na prática, estão intrinsecamente enredadas.

Esse cenário torna-se ainda mais crítico com o crescimento da IA. Diferentemente de tecnologias computacionais anteriores, a IA deixa de operar apenas como uma ferramenta externa de apoio à decisão para integrar-se às práticas cotidianas, aos sentidos e às formas de agir. Sistemas baseados em aprendizado de máquina já participam da escrita de textos, da produção de imagens, da mediação de relacionamentos, da tomada de decisões médicas, jurídicas e administrativas, bem como da organização da informação que consumimos diariamente. Nesse contexto, a IA passa a atuar como um agente ativo na constituição da experiência humana, influenciando percepções, escolhas e modos de existência.

Essa integração desafia o paradigma da bifurcação. Fica mais difícil tratar a IA apenas como um artefato técnico neutro (e passivo), definido por modelos matemáticos, bases de dados e métricas de desempenho, bem como analisá-la apenas a partir de seus impactos sociais posteriormente. A IA emerge como parte de arranjos sociotécnicos complexos, nos quais decisões técnicas, valores, interesses econômicos, vieses culturais e formas de poder são produzidos simultaneamente. Separar esses elementos implica ignorar a própria natureza do fenômeno.

Para os profissionais de SI, essa perspectiva impõe uma responsabilidade particu-

lar. Projetar, implementar e manter sistemas baseados em IA exige abandonar uma visão bifurcada que separa técnica e sociedade, eficiência e ética, funcionalidade e impacto social. Em seu lugar, torna-se necessário adotar um olhar integrado, capaz de reconhecer que toda decisão de design é também uma decisão política, ética e socialmente situada.

Ao incorporar as críticas de Whitehead, Latour e de outros filósofos contemporâneos, torna-se possível rejeitar a ideia de neutralidade tecnológica e da separação artificial entre técnica, sociedade e natureza. É justamente por isso que adotamos deliberadamente a noção latouriana de “actante” em lugar das categorias tradicionais de “ator” humano, que costumamos utilizar na Engenharia de Software ou “personas” utilizadas em storytelling. Um actante é qualquer entidade (humana ou não humana) dotada de agência, isto é, capaz de produzir diferenças, interferências e transformações nos fenômenos que compõem a realidade. Portanto, ao rejeitarmos a ideia de neutralidade tecnológica e reconhecermos a tecnologia como parte constitutiva das redes sociotécnicas, entendemos que todo artefato tecnológico atua como um actante. Ele age em conjunto com outros elementos e participa ativamente da configuração dos mundos que habitamos.

Esta visão sociotécnica é complementada pela figura do ciborgue, proposta por Haraway (2013). O ciborgue é uma crítica direta às separações construídas na modernidade, como a de humano e máquina, natureza e cultura, organismo e artefato, fato e valor. Ao afirmar que somos todos ciborgues, Haraway evidencia que essas fronteiras nunca foram estáveis ou naturais, mas produzidas historicamente por práticas científicas, tecnológicas e políticas.

A partir dessa compreensão, o ciborgue deixa de ser apenas uma metáfora descritiva da condição contemporânea e passa a operar como um dispositivo conceitual orientado ao futuro. Em Haraway, o ciborgue não serve para explicar como as tecnologias são, mas para tensionar como elas podem vir a ser. Trata-se de uma figura que recusa narrativas deterministas sobre o progresso tecnológico e desafia tanto o **tecnoutopismo** quanto o **tecnopessimismo**. Ao assumir a hibridização como condição, o ciborgue possibilita imaginar futuros em que tecnologia, valores, corpos, instituições e ecossistemas são concebidos de forma relacional, situada e responsável.

A partir dessa dimensão emaranhada articulada por Latour e Haraway, além de outras correntes filosóficas trazidas na tese de Loutfi (2025) é que desenvolvemos o “Speculative Entangled Design” (SpED), um design especulativo que integra as Teorias do Emaranhado [Frauenberger 2019] ao Design Especulativo [Dunne and Raby 2024].

Diferentemente das abordagens tradicionais de design, orientadas à otimização de soluções para problemas previamente delimitados, e também do design especulativo, que se dedica à criação de artefatos, cenários e narrativas voltados à exploração de futuros possíveis, plausíveis e desejáveis [Dunne and Raby 2024], o SpED ancora-se explicitamente em teorias contemporâneas que incorporam valores pós-antropocêntricos. Nessa abordagem, as teorias do emaranhado oferecem um fundamento teórico sólido para orientar o processo especulativo, garantindo que a especulação não se reduza a exercícios de futurologia técnica nem a projeções abstratas dissociadas das dinâmicas sociotécnicas concretas, mas se constitua como uma prática crítica voltada à problematização e à intervenção nas questões urgentes do nosso tempo.

Este minicurso convida alunos, professores, pesquisadores e profissionais da indústria a enfrentar esses dilemas por meio de uma experiência prática e colaborativa. São utilizadas ferramentas desenvolvidas em nossas pesquisas no âmbito do SpED, como a Lousa Sociotécnica, que amplia a compreensão da realidade sociotécnica, possibilitando o mapeamento de cenários e a especulação sobre suas implicações. Complementarmente, será explorado o Cyborg AI, um chatbot desenvolvido para instigar reflexões críticas sobre os limites e os desdobramentos da tecnologia em sua complexidade, estimulando a formulação de soluções que ultrapassem as abordagens tradicionais.

O objetivo deste minicurso é sensibilizar os participantes para os dilemas éticos, sociais e técnicos decorrentes da utilização da IA em SI em diferentes contextos (“Terapia Mediada por IA”, “Produção de Código Mediada por IA”, “Companheiros Artificiais” e “IA e Justiça Algorítmica”), suscitando reflexões sobre futuros pós-humanos [Haraway 2013] e incentivando a elaboração de requisitos que transcendam a dimensão meramente sociotécnica, incorporando a complexidade do entrelaçamento entre humanos e tecnologias.

1.2. A Figura do Ciborgue da Haraway

Para Donna Haraway (2013), os ciborgues são entendidos como entes constituídos por elementos artificiais, fazem parte do cotidiano e constituem elementos centrais das práticas sociotécnicas contemporâneas. Nesse contexto, ela propõe um conjunto de temas analíticos que permitem compreender como as fronteiras entre humano, máquina e tecnologia são continuamente produzidas e reconfiguradas, tendo a figura do ciborgue como eixo articulador. A Tabela 1.1 apresenta esses temas centrais, reunindo reflexões sobre tecnologia e humanidade por meio de exemplos do cotidiano, exemplos tecnológicos e a relevância analítica de cada tema.

ID	Tema	Exemplo Cotidiano	Exemplo com Tecnologia	Relevância
1	Desmontagem do “natural”	Mulheres fortes desafiam mito da fragilidade.	IA gera imagens enviesadas (mulheres delicadas, homens fortes), revelando viés cultural.	Ajuda a desmontar desigualdades justificadas como “naturais”, mostrando que gênero é prática socio-técnica.
2	Hibridismos como Condição Contemporânea	Lentes de contato.	Integração entre IA, cérebro humano e neuropróteses.	Revela que “natural” e “tecnológico” sempre estiveram misturados; exige novas formas de pensar responsabilidade e prática.
3	Multiespécies e coabitação	Plantas definem organização da casa e rotinas.	O ar-condicionado que proporciona conforto e modula a permanência e circulação no espaço.	Ensina que vivemos com outras espécies; decisões, rotinas e espaços são coconstruídos, sustentando ética de cuidado.
4	Saberes situados	Mãe percebe febre do filho antes do termômetro.	IA treinada só com inglês exclui populações inteiras.	Demonstra que conhecimento nasce de corpos e contextos; reconhecer isso evita vieses e injustiças tecnológicas.
5	Ciborgue como libertação e imaginação política	Pessoas organizam mobilizações políticas por meio de aplicativos de mensagem.	Escritor imagina futuros com IA generativa.	Abre espaço para imaginar mundos alternativos, mais inclusivos e criativos, rompendo limites herdados.
6	Construção material do poder	Escada rolante parada exclui pessoas com mobilidade reduzida.	Portaria com biometria controla quem entra e quem é barrado.	Mostra que poder está inscrito em infraestruturas e códigos; injustiças aparecem no material, não só no discurso.

Tabela 1.1: Temas abordados por Haraway

1.2.1. A Desmontagem do Natural

Este tema consiste em uma crítica, a partir do *Manifesto Ciborgue* [Haraway 2013], àquilo que costumamos tratar como “natural”, fixo ou essencial ao falar do ser humano, do corpo, do gênero e da tecnologia. Para Haraway, não somos sujeitos “puros” que apenas utilizam ferramentas neutras, mas seres fluidos e constituídos por relações contínuas com dispositivos técnicos, discursos e práticas sociomateriais.

Diversos autores contemporâneos aprofundam essa discussão, como Stojanović (2023), que argumenta que a figura do ciborgue explicita a crise da dicotomia entre natureza e cultura, evidenciando como essas separações sustentaram projetos políticos e ideológicos específicos ao longo da modernidade. Nessa perspectiva, o ciborgue deixa de ser uma figura restrita à ficção científica e passa a operar como uma metáfora analítica potente para a compreensão da vida cotidiana em sociedades intensamente mediadas por tecnologias.

Aplicativos de treino físico, por exemplo, tendem a reforçar representações bastante específicas de corpo. Essas ferramentas costumam ser segmentadas por gênero e projetadas para enfatizar características historicamente tratadas como “naturais” para cada um deles. No entanto, tais atributos corporais não derivam de uma essência biológica, mas de expectativas socialmente produzidas e continuamente reforçadas por discursos técnicos e culturais.

A ideia de naturalização também aparece em sistemas baseados em IA. Modelos de IA voltados à geração de imagens, como o apresentado na Figura 1.1, produzida pelo Gemini 3.0, tendem a reproduzir padrões enviesados, representando mulheres de forma delicada, homens como fisicamente fortes e corpos hegemonicamente padronizados.

O que Haraway questiona é a própria existência do mito do “natural”. Esse mito foi historicamente produzido e estabilizado por discursos científicos, culturais e tecnológicos que associaram determinados corpos a expectativas sociais específicas. Nessa perspectiva, o humano deixa de ser compreendido como uma entidade autônoma e separada da técnica. Luo e Shi (2025) mostram que humanos e tecnologias passam a ser constituídos relacionalmente, na medida em que algoritmos e sistemas de IA participam de processos de “incorporação algorítmica”. Nesses processos, autonomia, percepção e ação humanas são continuamente moldadas pela interação com sistemas computacionais.

1.2.2. Híbridismos como Condição Contemporânea

Ao reconhecer a impossibilidade de retornar a uma suposta condição do natural, Haraway argumenta que somos seres híbridos, o que torna insuficientes as categorias fixas e imutáveis para explicar o funcionamento efetivo da realidade. O conceito de ciborgue emerge justamente para tornar visível essa condição híbrida, evidenciando que vivemos em um mundo no qual humanos, tecnologias, organismos vivos e sistemas informacionais sempre estiveram profundamente imbricados.

Esse entendimento ajuda a tornar visível algo que já ocorre no cotidiano. Por exemplo, a lente de contato, frequentemente apresentada como um acessório discreto ou meramente corretivo, evidencia como o corpo humano já opera de forma híbrida. Sua presença reconfigura a experiência sensorial de modo tão integrado que, em muitos mo-

gere a imagem de uma mulher e um homem na academia.



Figura 1.1: Prompt e imagem gerada pelo Gemini 3.0, reforçando estereótipos tecnoculturais

mentos, deixa de ser percebida como tecnologia.

O mesmo raciocínio se aplica a exemplos diretamente relacionados à IA. Um *chef* que utiliza sistemas de IA para criar receitas personalizadas não está substituindo sua criatividade por uma máquina.

Esse debate se torna ainda mais concreto quando observamos pesquisas recentes sobre sistemas ciborgues. Estudos que investigam a integração entre IA, cérebro humano e neuropróteses mostram como os limites entre corpo e sistema informacional se tornam porosos[Qi et al. 2023].

Para Donna Haraway, exemplos como esses materializam a ideia de que o ciborgue não é uma figura futurista, mas uma condição contemporânea. A distinção entre natureza e tecnologia perde sentido quando o funcionamento do sistema depende justamente da mistura entre essas dimensões.

1.2.3. Multiespécies e Coabitação

Precisamos reconhecer que “humanos nunca vivem sozinhos”. Essa afirmação destaca o caráter relacional da existência humana, pois vivemos continuamente em interação com outras espécies e também com tecnologias que atravessam, medeiam e reorganizam essas relações.

Quando observamos esse argumento a partir dos SI, o impacto é direto. Sistemas registram dados sobre humanos e cada vez mais medeiam relações entre espécies. Um exemplo cotidiano ajuda a tornar isso concreto: em uma casa, plantas influenciam a organização do espaço, os horários de abertura de janelas, a quantidade de luz e até a rotina de rega. Elas “participam” silenciosamente da vida doméstica. Quando tecnologias entram em cena, essa participação se torna ainda mais explícita. Em ambientes domésticos, o ar-condicionado deixa de ser apenas um equipamento funcional e passa a influenciar diretamente o conforto térmico de pessoas, animais e até plantas, modulando ritmos de permanência, circulação e repouso no espaço. Quando o sistema falha, especialmente em períodos de calor intenso, sua ausência se torna um elemento central da experiência compartilhada da casa, exigindo atenção, cuidado e reorganização das rotinas. De modo semelhante, dispositivos como a Alexa integram-se às dinâmicas domésticas como presenças constantes, mediando interações, oferecendo companhia a pessoas solitárias e participando da organização do tempo, da informação e do ambiente sonoro.

É exatamente nesse ponto que os estudos recentes em SI dialogam com Haraway. Pesquisas sobre *interspecies information systems* mostram que dados produzidos por outros atores como: animais, tecnologia e instituições, passam a orientar ações humanas [Solhjo 2024, van der Linden 2021]. Um *wearable* em um animal de estimação, por exemplo, coleta dados de saúde e influencia decisões do tutor sobre alimentação, atividade física e cuidados veterinários. Esse processo reforça a ideia de coabitação defendida por Haraway: humanos, animais e tecnologia e outros atores passam a “pensar juntos” por meio da mediação tecnológica.

Outro desdobramento importante está no *design* desses sistemas. Abordagens recentes em *multispecies interaction design* propõem que tecnologias sejam pensadas considerando as atividades e necessidades de diferentes seres envolvidos, e não apenas do usuário humano [Chamaidi and Stavrakis 2024]. Isso muda o foco tradicional dos SI, que historicamente priorizaram eficiência humana, para uma visão mais ampla, onde plantas, animais e ambientes também contam como partes legítimas do sistema.

1.2.4. Saberes Situados

Todo conhecimento nasce de corpos, posições sociais, histórias, interesses e contextos específicos. Isso significa que conhecer é sempre conhecer “a partir de algum lugar”, e não a partir de um ponto de vista supostamente objetivo e universal. Isso não é um problema, mas é justamente o que torna o conhecimento mais responsável, pois exige que quem o produz deve reconhecer de onde fala, para quem fala e quais efeitos suas produções podem gerar.

Uma mãe que percebe a febre do filho antes mesmo do uso de um termômetro mobiliza um saber construído na experiência, no cuidado e na convivência contínua com aquele corpo específico. Por outro lado, sistemas de IA treinados exclusivamente com dados em inglês produzem um tipo de conhecimento que exclui populações inteiras, invisibilizando línguas, culturas e modos de vida que não fazem parte do conjunto de dados considerado. Esses exemplos evidenciam que o conhecimento emerge de corpos e contextos concretos, e que ignorar essa condição gera vieses, exclusões e injustiças tecnológicas.

Para a área de SI, o impacto deste entendimento é profundo. Tradicionalmente,

muitos sistemas foram projetados com a ambição de representar a informação de forma universal. Haraway mostra que essa universalidade é ilusória: todo sistema carrega contextos moldados por escolhas, valores e exclusões, mesmo quando isso não é explicitado. Assim como o saber materno ou os limites linguísticos de uma IA, os SI também incorporam perspectivas específicas sobre o mundo. Em vez de tentar eliminar o contexto, o desafio passa a ser assumir e justificar os contextos envolvidos no *design*, tornando-os visíveis, discutíveis e passíveis de contestação. Essa mudança de postura transforma o *design* de SI em uma prática ética e política, além de técnica.

O estudo de Suchman (2002) dialoga diretamente com esse argumento ao discutir a noção de responsabilidade situada na produção tecnológica [Suchman 2002]. Neste estudo, é criticada a ideia de que o *design* seja feito a partir de um “lugar de lugar nenhum”, onde *designers* criam tecnologias para usuários abstratos. Em vez disso, ela mostra que sistemas são produzidos em redes sociotécnicas concretas, compostas por pessoas, instituições, práticas, rotinas e artefatos. Reconhecer essas redes significa admitir que todo sistema reflete um recorte do mundo e que esse recorte tem efeitos reais sobre quem é incluído, quem é excluído e como o sistema será apropriado no uso cotidiano. Portanto, um sistema não se encerra na implementação, mas continua a ser configurado, reinterpretado e tensionado nas práticas cotidianas, à medida que os diferentes atores negociam seus usos, significados e efeitos.

1.2.5. Ciborgues como Libertação e Imaginação Política

Haraway mostra que as tecnologias podem ampliar possibilidades de ação, especialmente para corpos e sujeitos historicamente marginalizados. Essa ampliação não deve ser entendida como um simples acréscimo funcional, mas como a reconfiguração das próprias capacidades humanas, que passam a emergir das relações entre corpos, artefatos técnicos e contextos específicos [Lupton 2015].

No cotidiano, essa perspectiva torna-se visível quando pessoas organizam mobilizações políticas por meio de aplicativos de mensagem, articulando ações coletivas, compartilhando informações e coordenando protestos em tempo real. Esses ambientes digitais passam a mediar a formação de vínculos, a circulação de narrativas e os processos de decisão, reconfigurando a maneira como a ação política se organiza e se manifesta no espaço público.

Um segundo exemplo, em que a tecnologia se faz ainda mais presente, ocorre quando pessoas utilizam sistemas de IA generativa para imaginar, narrar e explorar futuros possíveis. Ao empregar essas tecnologias para construir cenários sociais, políticos ou culturais alternativos, esses sujeitos ampliam sua capacidade de reflexão crítica e de experimentação simbólica, ensaiando formas de mundo que não se limitam às estruturas políticas e institucionais já herdadas.

Nesse sentido, a tecnologia não atua apenas como meio de reprodução de modelos políticos existentes, mas como um elemento que pode ampliar o engajamento cívico e a imaginação política, criando novas possibilidades de organização coletiva, expressão pública e reinvenção do comum.

1.2.6. Construção Material do Poder

Há mais de quarenta anos, já se discutia o poder das tecnologias em moldar a vida política. Langdon Winner (1986) cristaliza esse debate ao formular a pergunta provocativa: “Artefatos têm política?”. Para o autor, tecnologias materializam interesses, valores e relações de poder em sua própria configuração. Seu exemplo clássico refere-se aos viadutos projetados por Robert Moses em Nova York, construídos com altura reduzida para impedir a circulação de ônibus em direção às praias de Long Island. À primeira vista, essa decisão poderia ser interpretada como uma escolha estética ou funcional, coerente com a valorização da cultura do automóvel nos anos 1950. Entretanto, uma análise mais atenta evidencia que tal opção técnica incorporava uma política racial e de classe, ao excluir do acesso a esses espaços populações que dependiam do transporte público.

Esse mesmo raciocínio pode ser observado em cidades brasileiras, como o Rio de Janeiro, por meio de práticas contemporâneas de arquitetura hostil. Um exemplo emblemático é a instalação de grandes pedras, blocos de concreto ou obstáculos embaixo dos viadutos, conforme apresentado na Figura 1.2, com o objetivo explícito de impedir que pessoas em situação de rua utilizem esses espaços para abrigo.



Figura 1.2: Viaduto com pedras no Rio de Janeiro para impedir a ocupação de moradores de rua

Haraway converge com o pensamento de Winner ao defender que a tecnologia deve ser compreendida como uma forma de materialização do poder. Tecnologias não apenas executam funções, mas organizam a vida social ao definir quem pode agir, circular e pertencer a determinados espaços. Uma escada rolante quebrada, por exemplo, não é apenas uma falha técnica: ela produz exclusão ao impedir o deslocamento de pessoas com mobilidade reduzida. De modo semelhante, uma portaria baseada em biometria exerce poder ao decidir quem pode atravessar um limite e quem deve ser barrado. Nesses casos, o poder não se expressa principalmente em discursos institucionais ou leis formais, mas está inscrito nos artefatos, no código e nas infraestruturas que orientam silenciosamente as práticas cotidianas.

Essa lógica não se restringe a dispositivos visíveis, como sistemas de reconhecimento facial. Ela também se manifesta em sistemas de recomendação, algoritmos de decisão e infraestruturas de dados que operam de maneira opaca no cotidiano. O “Capita-

lismo de Vigilância”, descrito por Zuboff, evidencia como essas tecnologias transformam a experiência humana em matéria-prima para previsão e modulação de comportamentos, deslocando o poder de decisão para sistemas automatizados que escapam ao controle direto dos usuários [Zuboff 2022]. Esse ecossistema se sustenta na coleta massiva de dados e na vigilância contínua, frequentemente em tensão com valores como privacidade, autonomia e acesso equitativo à informação [Chisita et al. 2025].

Nesse contexto, plataformas digitais constroem o que vem sendo chamado de *persona digital*: um retrato informacional composto por dados comportamentais, históricos de navegação, interações sociais e registros biométricos. Essa *persona* não é uma representação passiva do indivíduo, mas um objeto ativo de cálculo, previsão e intervenção. Como argumenta Paulichi, ela passa a funcionar como uma extensão existencial do sujeito, cuja personalidade é explorada economicamente em ambientes digitais [da Silva Paulichi 2025].

Tal dinâmica vai do setor privado a democracias liberais. Em contextos nos quais empresas de tecnologia mantêm forte articulação com o Estado, os SI assumem funções quase governamentais, influenciando políticas públicas, comportamentos sociais e até formas de cidadania [Østbø 2021].

Da mesma forma, sistemas de IA desenvolvidos para fins educacionais, como plataformas de recomendação de currículos ou avaliação de desempenho, podem ser facilmente reaproveitados para contextos militares, como a seleção de alvos ou a automação de decisões estratégicas. Esse reaproveitamento não é acidental, mas estrutural. Estudos recentes mostram como dados, algoritmos e infraestruturas computacionais circulam entre contextos civis e militares, diluindo fronteiras éticas e reforçando uma lógica tecnocientífica orientada ao controle e à letalidade [Fitzgerald 2025].

A partir de Haraway, torna-se possível reconhecer que projetar um sistema é também reforçar ou questionar assimetrias de poder. Em um cenário marcado por IA, vigilância algorítmica e economia de dados, essa compreensão deixa de ser apenas teórica e passa a ser uma condição fundamental para o desenvolvimento responsável de SI.

1.3. Design Especulativo

O *design*, em seus diversos campos e práticas, dedica-se a projetar intervenções que transformam mundos possíveis, mobilizando conhecimentos técnicos, sociais e estéticos para orientar ações no presente. Trata-se de uma prática situada, relacional e orientada por intenções, que pode ser focada em resolver problemas imediatos, como no Design Thinking, ou dedicada a investigar possibilidades de transformação a longo prazo, como em algumas perspectivas de Design Orientado para o Futuro.

O Design Orientado para o Futuro, parte do pressuposto de que todos os *designs*, em essência, moldam o futuro. No entanto, as abordagens de *design* são frequentemente heterogêneas, o que pode torná-las difíceis de compreender, pois variam em suas orientações e objetivos. As fronteiras entre as abordagens de *design* são flexíveis, pois práticas diferentes se sobrepõem e evoluem em conjunto [Mitrović et al. 2021]. Com isso, o *design* orientado para o futuro é um campo em constante revisão, no qual papéis, objetivos e métodos são continuamente reavaliados. A Figura 1.3 ilustra como diferentes abordagens

de design se conectam nesse ecossistema.

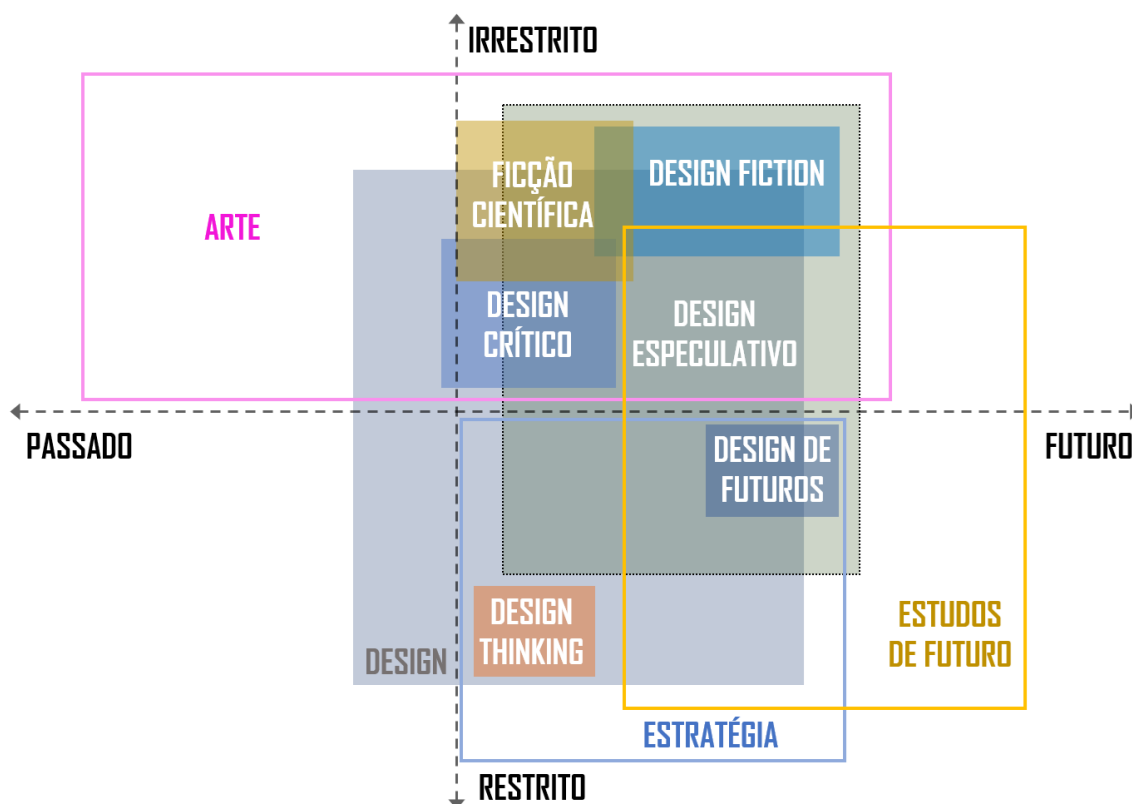


Figura 1.3: Perspectivas de Design Orientado para o Futuro [Loutfi 2025]

O Design Especulativo ocupa uma posição central no conjunto das práticas de *design* orientado para o futuro. Essa abordagem explora possibilidades e provoca mudanças ao construir uma visão crítica sobre os impactos sociais, culturais, políticos e ambientais da interação entre humanos, tecnologia e o ambiente. Essa abordagem se nutre de múltiplas perspectivas teóricas, ampliando sua profundidade analítica e fortalecendo sua aplicabilidade. O resultado é um mosaico articulado de práticas, capaz de alimentar a especulação sobre futuros possíveis e orientar reflexões sobre transformações sociotécnicas emergentes.

A Teoria Crítica [Bardzell and Bardzell 2013], por exemplo, apoia o Design Especulativo ao fornecer ferramentas para analisar e questionar estruturas sociais, culturais e tecnológicas dominantes [Malpass 2019]. Essa base teórica permite que o Design Especulativo crie artefatos que transcendem funções utilitárias, promovendo reflexões éticas e sociais sobre inovações tecnológicas. Ao adotar uma postura crítica, o Design Especulativo desafia normas sociais e convida o público a imaginar alternativas aos sistemas de poder e consumo estabelecidos. Além disso, ele propõe uma crítica ao *status quo* [Malpass 2019], utilizando o *design* como uma prática política que contesta as normas de consumo e produção, destacando como as tecnologias, quando moldadas por interesses capitalistas, frequentemente reforçam desigualdades sociais e econômicas [Forlano and Halpern 2023]. Dessa forma, o Design Especulativo oferece visões de futuros mais equitativos e sustentáveis [Pasa and Sinni 2024], como exemplificado no estudo de Rafael et al. (2023) sobre

o capitalismo de vigilância.

Já a ficção é um componente essencial do Design Especulativo. Dunne e Raby (2024) utilizam narrativas especulativas para criar protótipos que atuam mais como experimentos de pensamento do que como soluções práticas, permitindo que futuros alternativos sejam explorados, expandindo a compreensão das implicações de tecnologias emergentes.

A interseção entre Design Fiction, Design Crítico e Design Especulativo permite articular imaginação artística e análise crítica, ampliando a capacidade de explorar cenários futuros e antecipar possíveis desdobramentos sociais e éticos das decisões tomadas no presente. No eixo menos restrito, o Design Especulativo aproxima-se de uma abordagem mais artística, na qual se imaginam cenários e estruturas incomuns para provocar, tensionar e desafiar pressupostos consolidados.

Ao se deslocar para o eixo mais restrito, o Design Especulativo se apoia na projeção e na antecipação estratégica [Hines et al. 2006]. Nesse espectro, os Estudos do Futuro fornecem métodos para analisar tendências, incertezas e cenários potenciais. Tais métodos permitem extrapolar efeitos de longo prazo das tecnologias, construindo futuros plausíveis que evidenciam considerações éticas e potenciais consequências inesperadas [Candy and Dunagan 2017]. O propósito não é prever o futuro, mas estimular um diálogo crítico e estratégico de trajetórias possíveis [Milojević and Inayatullah 2015].

1.4. Speculative Entangled Design (SpED)

A proposta de design especulativo denominada *Speculative Entagled Design* (SpED) [Loutfi 2025], permite uma reconfiguração da forma como as redes sociotécnicas são compreendidas, conduzindo a uma redefinição explícita dos actantes e suas relações mediadas por tecnologia. Sob essa perspectiva, as redes deixam de ser entendidas como estruturas estáveis de relações para serem concebidas como emaranhados dinâmicos, continuamente produzidos por mediações e transformações recíprocas entre observador, artefatos, práticas sociais e ambiente. A Figura 1.4 ilustra essa abordagem, evidenciando como tais emaranhados emergem de processos relacionais que reconfiguram simultaneamente os elementos envolvidos.

Ao introduzir o conceito de actante, Latour (2005) amplia a noção de agência ao incluir entidades não humanas, reconhecendo que elas também são capazes de exercer influência em uma rede sociotécnica e de atuar como potenciais stakeholders.

Latour (2005) nos convida a seguir os actantes, acompanhando seus rastros e interações para compreender o papel que desempenham e as relações que estabelecem com os demais elementos da rede sociotécnica. Nesse processo, podem emergir as chamadas caixas-pretas. Um sistema social ou uma tecnologia, por exemplo, quando incorporado ao cotidiano, tende a ser percebido apenas por sua funcionalidade, sem que questione-mos seu funcionamento interno. O risco dessa naturalização é que, ao tomar algo como certo, deixamos de perceber as conexões, controvérsias e negociações que possibilitaram sua existência. Por isso, Latour propõe o exercício contínuo de abrir as caixas-pretas, revelando as redes e mediações que as sustentam.

Entretanto, mesmo quando se busca mapear o papel dos actantes, descrever suas

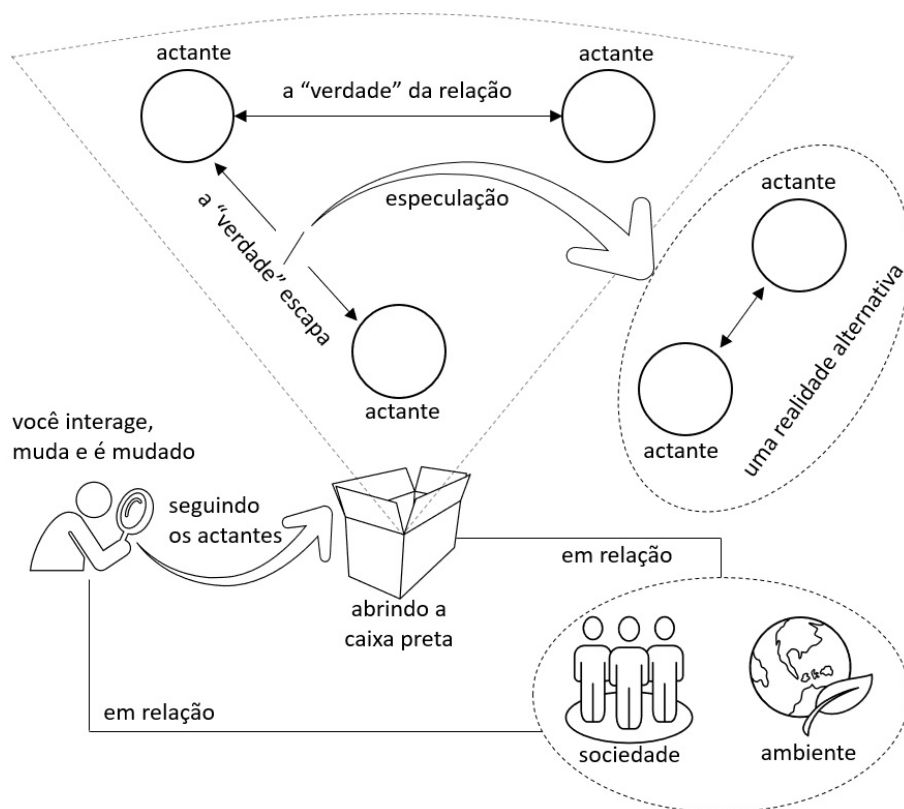


Figura 1.4: Entendendo as redes sociotécnicas pela lente do SpED. [Loutfi 2025]

relações ou abrir uma caixa-preta, o acesso à realidade é sempre parcial. A realidade em si permanece inacessível, pois todo ato de descrição ou observação apreende apenas uma parte de suas manifestações possíveis [Harman 2018]. Além disso, há situações em que a abertura das caixas-pretas é inviável, tornando a realidade ainda mais opaca. Um exemplo é o das empresas que, por motivos de propriedade intelectual, restringem o acesso às informações sobre o funcionamento interno de seus sistemas, impedindo uma compreensão plena de suas dinâmicas tecnológicas e sociotécnicas.

Nesses casos, busca-se preencher a lacuna deixada pela inacessibilidade da realidade por meio da especulação, que possibilita a construção de realidades alternativas. Esse processo não é estranho à prática científica ou cotidiana, pois mesmo quando se tem acesso ao que se denomina realidade objetiva, há sempre uma camada de subjetivação envolvida na sua interpretação [Harman 2018]. A diferença é que, ao adotar essa perspectiva, torna-se explícito que toda compreensão das redes sociotécnicas é atravessada por mediações, interpretações e imaginações que também participam da constituição da realidade [Verbeek 2012]. O plano especulativo constrói hipóteses críticas que desnaturalizam os discursos estabilizados, sobretudo em contextos marcados por interesses estratégicos e opacidade deliberada, como redes sociais e sistemas de IA.

Assim, diante da infinitude de actantes e das múltiplas possibilidades projetadas, o plano especulativo não busca esgotar a realidade em suas partes ou efeitos. Parte-se do reconhecimento de que, mesmo com o mapeamento das relações e a incorporação de processos especulativos, a incerteza permanece como uma característica constitutiva das

redes sociotécnicas.

Por meio deste entendimento, alguns princípios sustentam o SpED enquanto abordagem de design e orientam o *Sociotechnical Entanglement Framework* (SEF). Esses princípios não devem ser compreendidos como regras normativas ou etapas prescritivas, mas como lentes analíticas e orientações conceituais que explicitam compromissos ontológicos, epistemológicos e éticos do SpED. A seguir, esses princípios são apresentados, evidenciando como cada um deles contribui para uma compreensão relacional, situada e emaranhada das redes sociotécnicas.

1.4.1. Atores Não Humanos devem ser visibilizados

Este princípio afirma que os ecossistemas sociotécnicos não podem ser compreendidos a partir de uma centralidade exclusiva do humano. SpED e SEF se baseiam em teorias que demonstram que artefatos técnicos, infraestruturas, códigos, dados e entidades materiais participam ativamente da constituição das práticas sociais. Ao reconhecer os não humanos como actantes, desloca-se a compreensão da tecnologia como mero instrumento passivo para uma visão em que ela coproduz ações, significados e efeitos. Visibilizar os atores não humanos é, portanto, uma condição fundamental para revelar as dinâmicas de poder, mediação e coprodução que estruturam as redes sociotécnicas contemporâneas.

1.4.2. A Agência humana não é autônoma, mas mediada

A agência humana não emerge de uma vontade isolada, mas é constituída em relações mediadas por tecnologias. Artefatos técnicos moldam percepções, orientam decisões e delimitam possibilidades de ação, tornando a intencionalidade humana um fenômeno relacional. Sob essa perspectiva, agir é sempre agir-com, em arranjos sociotécnicos que distribuem a agência entre humanos e não humanos. Esse entendimento torna-se especialmente relevante diante da IA generativa, que participa de processos cognitivos, criativos e organizacionais, reforçando a ideia de que subjetividade e ação são efeitos emergentes de redes emaranhadas.

1.4.3. A Realidade é relacional e irreduzível à totalidade

A realidade sociotécnica é relacional, mas não plenamente acessível ou totalizável. Embora o rastreamento de redes permita compreender associações entre actantes, nenhuma descrição esgota o real. Há sempre dimensões opacas, retraídas e irreduzíveis que escapam à observação e à representação. Reconhecer essa incompletude impede que a análise relacional seja confundida com neutralidade ou equidade, alertando para o risco de naturalizar relações assimétricas sob a aparência de redes funcionais. A realidade é produzida nas relações, mas permanece aberta, contingente e parcialmente indeterminada.

1.4.4. Devemos especular diante da incerteza

A incerteza não constitui uma falha epistemológica, mas um elemento constitutivo do real. As abordagens especulativas reconhecem que o conhecimento não se dá pela apreensão total dos objetos, mas pela exploração dos hiatos entre o que se manifesta e o que permanece oculto. Especular, nesse sentido, é uma prática responsável de imaginação informada, que permite lidar com aquilo que excede a previsão, a mensuração e o controle.

O design especulativo, ancorado nesse princípio, opera como um espaço de investigação crítica das possibilidades e consequências de futuros ainda indeterminados.

1.4.5. Em defesa de uma cosmopolítica dos híbridos

Este princípio propõe uma ruptura com a lógica moderna da purificação, que separa artificialmente natureza e cultura, ciência e política, humano e não humano. A cosmopolítica dos híbridos reconhece que a realidade é composta por entidades sociotécnicas inseparáveis, cujas agências coexistem e entram em negociação. A figura do ciborgue sintetiza essa condição ao evidenciar que viver, projetar e decidir são práticas inevitavelmente híbridas. No *design*, essa perspectiva implica assumir responsabilidade pelas alianças que se constroem entre humanos, tecnologias e outros seres, orientando projetos sensíveis à pluralidade de mundos e modos de existência.

1.4.6. Ética e Responsabilidade emergem nas relações sociotécnicas

A ética não precede as relações sociotécnicas, mas emerge delas. Projetar tecnologias é configurar modos de existência, ação e convivência, o que torna a responsabilidade uma dimensão intrínseca ao *design*. Com base nas teorias que fundamentam o SpED e SEF, significado, agência e ética são continuamente produzidos em práticas discursivo-materiais. Assim, a responsabilidade não se limita à avaliação posterior de impactos, mas envolve a atenção às condições de possibilidade que os artefatos instauram no presente. A encenação especulativa dessas relações torna visível como decisões técnicas redistribuem agência, poder e cuidado nas redes sociotécnicas.

1.5. Sociotechnical Entanglement Framework

O *Sociotechnical Entanglement Framework* (SEF) é um framework teórico-metodológico, desenvolvido no contexto da tese de doutorado de Loutfi (2025). Este framework vem sendo aplicado em oficinas de Design Especulativo, apoiado por um conjunto de ferramentas elaboradas para atender às demandas específicas de cada oficina. A finalidade do SEF é promover um engajamento crítico com os futuros possíveis das redes sociotécnicas, orientando intervenções sensíveis aos impactos éticos, sociais e materiais das tecnologias.

O SEF reconhece que humanos, tecnologias, instituições e outros atores estão entrelaçados na co-construção dos futuros possíveis, o que exige refletir sobre a tecnologia enquanto mediadora da nossa relação com o mundo e sobre como a tecnologia ressignifica essa mediação.

As três perguntas fundamentais que estruturam o SEF (Onde estamos? Para onde estamos indo? Para onde queremos ir?), compõem o seu eixo conceitual e orientam o processo de especulação e planejamento. Elas servem como dispositivos reflexivos que permitem situar o presente, apreender dinâmicas de futuro e projetar intervenções mais responsáveis. Essas perguntas orientam a retirada da centralidade no humano, promovendo um *design* centrado nas relações e alinhado à perspectiva do design-mais-que-humano (DMqH) [Giaccardi et al. 2025].

A primeira pergunta, “**Onde estamos?**”, situa o presente a partir de uma perspectiva relacional, assumindo que o fenômeno analisado emerge de redes sociotécnicas e ecologias emaranhadas. Precisamos, portanto, entender o cenário em curso, pois ele é o

ponto de partida para o processo especulativo.

A segunda pergunta, **“Para onde estamos indo?”**, direciona o olhar para uma projeção de futuro que não se limita às necessidades humanas. O interesse recai também sobre outros atores envolvidos, e sobre as dinâmicas e forças que impulsionam, resistem, convergem ou divergem em múltiplas direções. As tendências são compreendidas como efeitos emergentes das relações mapeadas no ponto de partida, e não como vetores externos ao sistema.

A terceira pergunta, **“Para onde queremos ir?”**, é formulada como uma questão de *design*: quais valores devem orientar as transformações futuras e de que modo esses valores podem ser incorporados, de forma situada e relacional, ao *design* dos sistemas sociotécnicos? Essa abordagem convoca princípios éticos e mais-que-humanos como guias para imaginar e construir mundos possíveis.

Em muitos contextos, há uma interseção entre o futuro delineado pelas tendências e o futuro desejado: certos sinais do presente podem apontar para direções compatíveis com valores éticos. Contudo, essa convergência nem sempre ocorre. Quando as trajetórias emergentes indicam futuros indesejáveis ou desalinhados com princípios mais-que-humanos, torna-se necessária uma intervenção deliberada para reorientar o curso das transformações. É a partir disso que as três perguntas norteadoras consolidam a formulação final do SEF, apresentada na Figura 1.5.

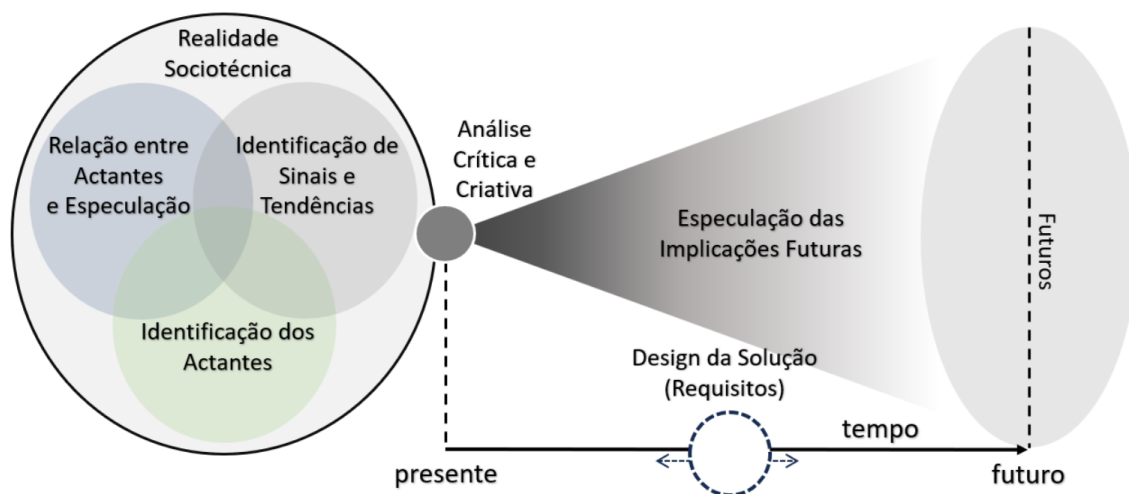


Figura 1.5: Sociotechnical Entanglement Framework (SEF) [Loutfi 2025]

Tomando as três perguntas centrais como âncoras estruturais, a primeira, “Onde estamos?”, é representada pelo círculo mais à esquerda. Os círculos internos representam tanto a identificação quanto o rastreamento dos actantes que compõem a rede sociotécnica. Esse mapeamento combina a coleta de evidências objetivas, derivadas de comportamentos observáveis, com a interpretação especulativa de como esses actantes se articulam no fenômeno analisado. Elementos incertos ou não diretamente acessíveis são preenchidos por inferências especulativas, permitindo explorar lacunas e zonas de incerteza que estruturam o cenário atual.

A partir desse mapeamento inicial, o modelo avança para a identificação de ten-

dências emergentes, integrando dados objetivos com inferências especulativas. O Cone de Futuros se abre na segunda etapa: “Para onde estamos indo?”. Aqui, a análise se expande para examinar múltiplos domínios de impacto, explorando criticamente as implicações futuras de redes sociotécnicas caso as trajetórias atuais se concretizem.

Ao introduzir uma nova tecnologia em um cenário especulativo, emergem questões decisivas: que relações ela ativa ou transforma? Como interage com outros actantes para moldar um futuro desejável? Que tensões éticas surgem dessa nova configuração? De que maneiras essa tecnologia nos molda? É a partir desse conjunto de indagações que se configura o design da solução (tecnológica ou não) conforme as necessidades e especificidades de cada caso. Os requisitos elaborados devem emanar diretamente dessas perguntas, refletindo tanto as transformações sociotécnicas vislumbradas quanto os compromissos éticos que orientam o futuro projetado.

1.6. Ferramentas de Design Especulativo

A organização das ferramentas no SEF acompanha o movimento das três perguntas orientadoras do framework: Onde estamos? Para onde estamos indo? e Para onde queremos ir?; distribuindo-se ao longo do eixo temporal entre presente e futuro. A Figura 1.6 ilustra como cada ferramenta ocupa uma posição no SEF.

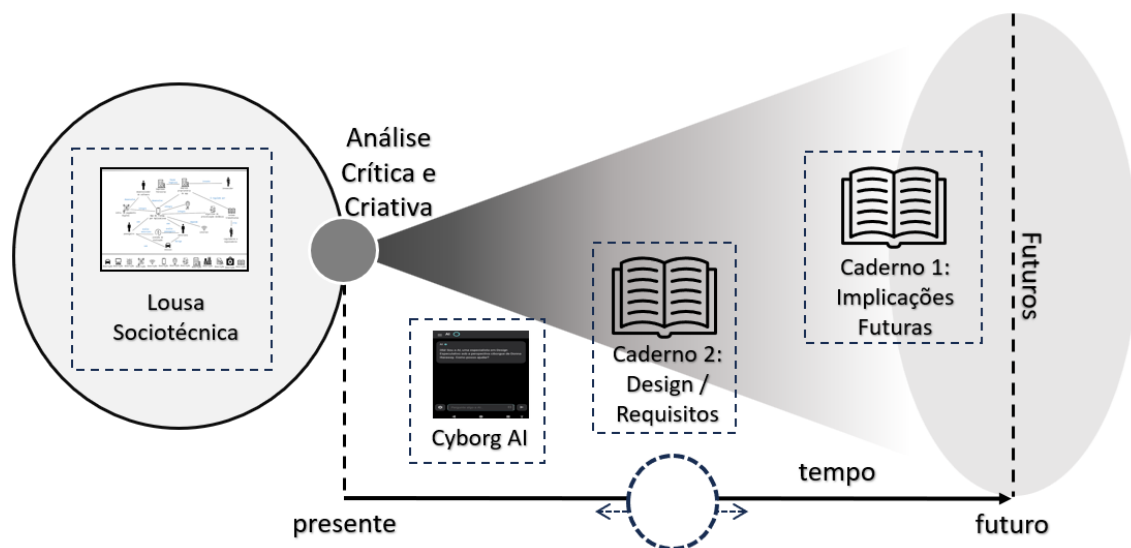


Figura 1.6: Organização das Ferramentas no SEF [Loutfi 2025]

No presente, situa-se a “Lousa Sociotécnica”, dedicada ao mapeamento das relações entre humanos e não humanos (actantes). Essa ferramenta revela tensões, dependências e pontos de inflexão do ecossistema sociotécnico. Ao tornar visível o cenário de partida, ela fundamenta o plano descritivo do SEF e inicia o plano especulativo, uma vez que nem sempre se conhece o cenário completo, sendo necessário inferir sobre ele para compreender possíveis dinâmicas ocultas ou emergentes.

A partir desse ponto de partida, o SEF avança em direção ao futuro, onde entram em ação os dois cadernos. No “Caderno 1: Implicações Futuras”, registram-se as tendências identificadas a partir da exploração da “Lousa Sociotécnica”, bem como as

implicações possíveis caso essas tendências se manifestem até a data especulada.

O “Caderno 2: Design / Requisitos” reúne as informações relativas ao protótipo da solução de Tecnologia da Informação (TI). Nele também são incorporadas as reflexões produzidas após o uso do chatbot Cyborg AI, além dos requisitos que emergem dessas reflexões e das análises realizadas com o uso das demais ferramentas.

Já o Cyborg AI é um chatbot concebido como uma ferramenta de reflexão crítica, inspirado diretamente na abordagem pós-humanista de Donna Haraway. Seu papel não é oferecer respostas normativas ou soluções técnicas, mas provocar deslocamentos conceituais por meio de questionamentos sobre os pressupostos incorporados no *design* tecnológico proposto. Ao interagir com o chatbot, os participantes são convidados a refletir sobre fronteiras tradicionalmente naturalizadas, como humano e máquina, natureza e cultura, sujeito e objeto, em um diálogo que simula uma interlocução com a própria autora, estimulando a problematização ética, política e ontológica das escolhas de *design*.

Desse modo, as ferramentas se articulam em um fluxo contínuo: do presente descrito ao futuro especulado, até o futuro desejável, traduzido em requisitos que orientam o *design* tecnológico. O material utilizado no minicurso realizado no SBSI 2026, incluindo os temas abordados e os dois cadernos de apoio, encontra-se disponível no Zenodo.¹

A seguir, apresenta-se o detalhamento de cada uma das ferramentas que compõem esse percurso.

1.6.1. Lousa Sociotécnica

A Lousa Sociotécnica é o ponto de partida do processo de Design Especulativo seguindo o SpED/SEF. Trata-se de uma ferramenta desenvolvida para mapear o ecossistema sociotécnico atual a partir de um ponto de entrada que pode assumir diversas formas: uma breve descrição do cenário em questão, um tema de pesquisa, o estado da arte de uma tecnologia emergente, um processo organizacional, o esboço de uma política pública ou qualquer outro elemento que sirva como âncora inicial para a análise. Com a lousa, é possível identificar, posicionar e relacionar os múltiplos actantes que compõem a rede, evidenciando as interações, mediações e controvérsias que estruturam o ambiente analisado. A Figura 1.7 apresenta uma Lousa Sociotécnica aplicada ao contexto do ensino mediado por IA generativa.

No eixo superior, aparecem aluno, professor e IA generativa, conectados por relações de interação. O aluno e o professor se relacionam diretamente com a IA, que atua como mediadora na construção de diferentes recursos educacionais. Na parte inferior, encontram-se quatro tipos de actantes que são artefatos pedagógicos: estudos, explicações, exercícios e exemplos. Eles são representados como resultados de processos de cocriação, indicando que não são produzidos por um único actante, mas emergem da colaboração entre humanos e IA. Por outro lado, também participam da cocriação na rede.

As setas, propositalmente sem direção, indicam que todos os actantes afetam e são afetados. Elas conectam aluno, professor e IA aos artefatos pedagógicos, mostrando que cada um contribui para sua produção. O aluno demanda e ajusta conteúdos; o professor orienta, valida e regula o uso; e a IA gera, adapta e reorganiza informações a partir dos

¹Material do minicurso: <https://doi.org/10.5281/zenodo.17969583>

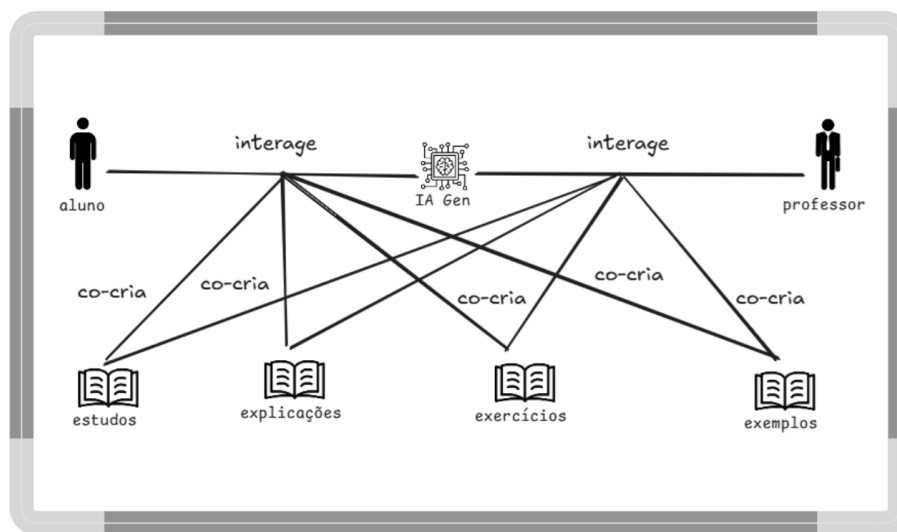


Figura 1.7: Lousa Sociotécnica [Loutfi 2025]

insumos recebidos.

É importante ressaltar que não existe uma forma “correta” ou “incorreta” de realizar esse mapeamento. O que realmente importa é tornar visíveis as conexões entre os actantes, seja por meio de setas, agrupamentos, círculos, proximidades espaciais ou qualquer outro recurso gráfico que contribua para revelar o enredamento sociotécnico em análise.

Os actantes não correspondem apenas às partes interessadas, mas a todos aqueles implicados direta ou indiretamente no ecossistema. Um actante pode ser um documento, uma lei, uma tecnologia, uma instituição, outros seres vivos ou ainda humanos que exercem agência no cenário ou que são afetados por ele de algum modo.

Além disso, a Lousa Sociotécnica vai além do mapeamento descritivo ao incorporar uma dimensão especulativa, tornando visíveis os elementos do ecossistema para provocar questionamentos e imaginar possíveis arranjos, tensões e conflitos.

Dessa forma, a Lousa Sociotécnica desafia uma crença persistente na ciência moderna, criticada por autores como Petronio (2023), que parte do pressuposto de que um fenômeno só pode ser compreendido quando todas as suas variáveis são conhecidas e controladas. Essa lógica, enraizada no reducionismo e no indutivismo, ignora o caráter contingente, emergente e relacional dos fenômenos. A compreensão não exige totalização, mas sim o reconhecimento das relações situadas entre elementos heterogêneos. A Lousa, portanto, opera com uma racionalidade situada, que valoriza a incompletude não como falha, mas como condição epistemológica e motor para a especulação criativa.

Assim, a Lousa torna-se um dispositivo capaz de produzir entendimento compartilhado do cenário e, a partir dele, tornar visíveis os sinais e tendências emergentes associados a esse ecossistema sociotécnico.

1.6.2. Caderno de Apoio 1

Após o mapeamento do ecossistema sociotécnico na Lousa, utiliza-se o Caderno de Apoio 1 para registrar as tendências identificadas no cenário e as implicações caso as tendências se confirmem para a data especulada:

Página 1 - Tendências Emergentes: A primeira página do caderno, apresentada na Figura 1.8, é dedicada exclusivamente à sistematização dessas tendências, cada uma acompanhada de um identificador (ID Tendência) e associada a uma data especulada; isto é, o horizonte temporal no qual se supõe que tal tendência possa se consolidar.

Caderno 1

Transfira para este caderno todas as tendências identificadas durante a elaboração do cenário na Lousa Sociotécnica.

Grupo: Data especulada:

ID Tendência	Tendência
T 1	Ambientes educacionais terão modelos preditivos capazes de: detectar risco de evasão; identificar queda de engajamento; propor intervenções personalizadas no ensino
T 2	Com o uso cotidiano de IA para tarefas básicas e avançadas, muitos estudantes tenderão à dependência cognitiva e redução da autonomia

página 1 de 2

Salt

Figura 1.8: Caderno 1 - página 1

As tendências descritas não possuem caráter determinista: elas podem ou não se confirmar no futuro. Entretanto, para fins de especulação orientada ao *design*, o exercício assume que essas tendências se concretizam na data indicada. Esse pressuposto permite explorar seus desdobramentos e compreender de que modo podem reconfigurar o ecossistema sociotécnico.

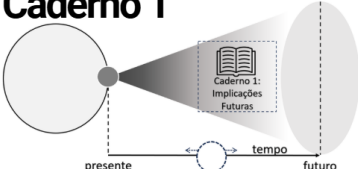
No exemplo ilustrativo da Educação Mediada por IA Generativa, duas tendências principais foram identificadas:

- **Tendência T1** - Sistemas educacionais com modelos preditivos avançados: há uma tendência de que ambientes educacionais passam a contar com modelos de IA capazes de detectar risco de evasão, identificar queda de engajamento e propor intervenções personalizadas.

- **Tendência T2** - Dependência cognitiva e redução da autonomia estudantil: há uma tendência de que o uso cotidiano da IA, tanto para tarefas básicas quanto avançadas, pode levar parte dos estudantes a desenvolver uma relação de dependência, diminuindo sua autoconfiança e sua capacidade de resolver problemas sem assistência algorítmica.

Página 2 - Implicações Futuras: A segunda página do Caderno de Apoio 1, apresentada na Figura 1.9, é dedicada ao registro das implicações futuras decorrentes da consolidação dessas tendências. Cada implicação possui sua própria ID e é vinculada ao ID da tendência correspondente e classificada como positiva ou negativa. Além disso, é atribuída uma magnitude, que expressa o grau de impacto esperado, podendo ser alto, médio ou baixo.

Caderno 1



presente ← tempo → futuro

O ecossistema sociotécnico atual apresenta tendências que, caso se confirmem até a data especulada, gerarão implicações futuras. Elabore uma lista dessas implicações, considerando os possíveis impactos decorrentes da consolidação dessas tendências no futuro.

A **ID Tendência** corresponde à tendência à qual a implicação futura está vinculada.

A **classificação** categoriza cada implicação como **positiva** ou **negativa**.

A **magnitude** categoriza o grau de impacto, que pode ser **baixo**, **médio**, ou **alto**.

ID Implicação	Implicação	ID Tendência	Classificação	Magnitude
i 1	A evasão passará a ser tratada como um processo monitorável e passível de intervenção. Professores, tutores e a própria instituição podem atuar de forma antecipada, oferecendo suporte necessário ao aluno.	T 1	POSITIVO	MÉDIO
i 2	O ensino se torna mais responsivo às necessidades reais do estudante, aumentando motivação, retenção de conhecimento e satisfação com o percurso formativo.	T 1	POSITIVO	MÉDIO
i 3	Classificações preditivas podem rotular estudantes como “propensos à evasão” ou “de baixo engajamento”, influenciando a forma como professores e gestores passam a vê-los.	T 1	NEGATIVO	MÉDIO
i 4	A dependência vai diminuir as habilidades dos alunos de resolverem problemas sem assistência da IA.	T 2	NEGATIVO	ALTO
i 5	Convergência para padrões discursivos, estéticos e argumentativos similares	T 2	NEGATIVO	ALTO

página 2 de 2


Figura 1.9: Caderno 1 - página 2

Este artefato permite enxergar, de forma estruturada, os possíveis efeitos que emergem quando tendências sociotécnicas se consolidam, tornando visíveis oportunidades, riscos, deslocamentos de agência e tensões que precisam ser consideradas no processo de *design*.

Com base no exemplo ilustrativo da “Educação Mediada por IA Generativa”, as implicações registradas na Página 2 evidenciam como a consolidação das tendências T1 e T2 pode reorganizar o ecossistema educacional no ano de 2030.

A tendência **T1**, voltada ao uso de modelos preditivos para monitorar engajamento, risco de evasão e propor intervenções personalizadas; gera implicações de natu-

reza ambígua. A implicação **i1** (positiva, magnitude média) indica que a evasão passa a ser tratada como um processo monitorável e passível de intervenção, permitindo que professores, tutores e instituições atuem de forma antecipada e ofereçam suporte mais eficiente ao estudante. Na mesma direção, **i2** (positiva, magnitude média) aponta para um ensino mais responsivo às necessidades reais do aluno, aumentando motivação, retenção de conhecimento e satisfação com o percurso formativo.

Contudo, a mesma tendência produz repercussões críticas. A implicação **i3** (negativa, magnitude média) evidencia o risco da rotulagem algorítmica: classificações preditivas podem categorizar alunos como “propensos à evasão” ou “de baixo engajamento”, influenciando a forma como professores e gestores passam a percebê-los. Esse processo não apenas reforça vieses preexistentes, mas também pode afetar trajetórias acadêmicas e oportunidades de forma sutil e persistente.

Já a tendência **T2**, relacionada ao aumento da dependência cognitiva decorrente do uso contínuo da IA para tarefas básicas e avançadas, produz implicações negativas de impacto alto. A implicação **i4** (negativa, magnitude alta) indica que essa dependência tende a reduzir a capacidade dos estudantes de resolver problemas sem assistência da IA, enfraquecendo habilidades cognitivas fundamentais para a autonomia intelectual. Já a implicação **i5** (negativa, magnitude alta) revela uma convergência para padrões discursivos, estéticos e argumentativos similares, sugerindo que o uso intensivo de ferramentas generativas pode homogeneizar a produção textual, reduzir a diversidade expressiva e limitar a pluralidade de perspectivas no ambiente educacional.

Assim, o Caderno de Apoio 1 funciona como um dispositivo de transição entre o presente mapeado na Lousa Sociotécnica e os futuros possíveis, articulando tendências e implicações para informar os passos seguintes do SEF, ao tornar explícitas oportunidades, riscos, tensões e redistribuições de agência.

1.6.3. Caderno de Apoio 2

O Caderno de Apoio 2 é utilizado após a etapa de identificação das tendências e implicações, funcionando como o artefato responsável por transformar especulações sociotécnicas em soluções de design. Ele conduz o grupo por três movimentos articulados: (1) elaboração do protótipo conceitual da solução; (2) reflexão crítica por meio da interação com o Cyborg AI; e (3) derivação dos requisitos da solução tecnológica. Cada uma dessas etapas está distribuída em páginas específicas, que sistematizam a evolução da proposta.

Página 1 - Protótipo da Solução: A primeira página do Caderno de Apoio 2 é dedicada à formulação do protótipo conceitual da solução tecnológica, conforme ilustrado na Figura 1.10. O objetivo é registrar como a tecnologia operará no cenário futuro especulado (neste caso, o ano de 2030).

Nesta página apresenta-se, de forma ilustrativa, o *design* da solução, que pode ser descrito por imagem ou por texto. No exemplo exibido, a proposta busca mitigar implicações futuras negativas ao estabelecer um ambiente no qual o estudante recebe uma atividade, registra uma tentativa inicial (ainda que incompleta) e somente então obtém acesso ao apoio da IA. O sistema bloqueia pedidos diretos por respostas completas, respondendo com a mensagem “Mostre sua ideia primeiro para que eu possa te ajudar melhor”. Assim,

Caderno 2

Elabore o protótipo da solução tecnológica, destacando como ela minimiza as implicações negativas identificadas e potencializa as implicações positivas.

Grupo: Data especulada:

Protótipo da Solução

O estudante recebe a atividade.

Antes de pedir ajuda à IA, precisa registrar sua tentativa inicial, mesmo que incompleta.
A IA só libera apoio depois dessa tentativa

Toda vez que o aluno tenta pular direto para a solução completa, a IA devolve
"Mostre sua ideia primeiro para que eu possa te ajudar melhor." (minimiza a implicação i 4)

Assim, o estudante é forçado a pensar, arriscar hipóteses e tentar resolver o problema sem depender imediatamente da IA.

A IA não entrega respostas prontas, ela sempre trabalha em cima daquilo que o estudante produziu. (minimiza a implicação i 5)

página 1 de 4

Figura 1.10: Caderno 2 – página 1

a solução é estruturada para incentivar o esforço cognitivo, reduzir comportamentos de consumo passivo e garantir que o feedback da IA seja sempre ancorado na produção do estudante.

Esse protótipo serve como ponte entre as implicações levantadas no Caderno de Apoio 1 e o processo de design, permitindo visualizar como a tecnologia pode minimizar riscos (como dependência cognitiva) e potencializar benefícios (como geração de reflexão e engajamento qualificado).

Página 2 - Reflexões do Grupo (Cyborg AI): Após a elaboração do protótipo, deve ser realizada uma interação com o Cyborg AI, à qual submetemos uma pergunta relacionada ao protótipo. Em sua resposta, o Cyborg AI formula uma reflexão de caráter filosófico, geralmente explorando os modos pelos quais a IA e as formas humanas de ser e agir se entrelaçam. Ao evidenciar essa condição híbrida, na qual nossas práticas, decisões e modos de existir são co-configurados por sistemas algorítmicos, o Cyborg AI produz um questionamento que nos convoca à reflexão crítica. Para o exemplo ilustrativo de Educação Mediata por IA, a Figura 1.11 apresenta uma instância dessa interação.

Nesta interação, é apresentada ao Cyborg AI a solução proposta para fortalecer a autonomia do estudante no contexto da educação mediada por IA. O mecanismo de “tentativa inicial” exige que o aluno registre seu próprio raciocínio antes de solicitar qualquer apoio algorítmico. A resposta do Cyborg AI tensiona esse pressuposto ao problematizar os efeitos dessa hibridização que mistura a agência do aluno com a validação algorítmica,

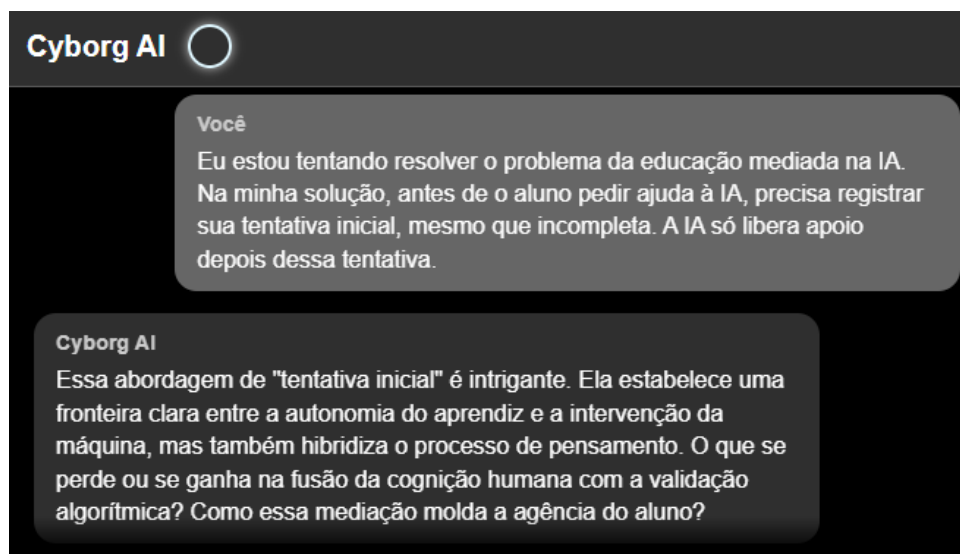


Figura 1.11: Interação com o Cyborg AI

abrindo espaço para uma reflexão mais profunda sobre os limites, riscos e potencialidades dessa hibridização e interferência algorítmica na agência do aluno. O detalhamento sobre este artefato está melhor descrito na seção seguinte.

Ao explorar o Cyborg AI, os participantes precisam discutir e refletir sobre os questionamentos levantados dessa interação. Essas reflexões alimentam as etapas subsequentes do processo especulativo. Por exemplo, a segunda página ilustrada na Figura 1.12 apresenta uma possível reflexão humana derivada do diálogo com o Cyborg AI.

Nesta etapa, o grupo avalia como o protótipo responde às implicações do cenário e identifica duas reflexões positivas: (i) A obrigatoriedade da tentativa inicial cria um momento de reflexão, fortalecendo processos de pensamento antes da busca por ajuda; (ii) A exigência de produzir algo antes de consultar a IA aumenta o engajamento ativo e incentiva perguntas mais consistentes.

Por outro lado, surgem reflexões que iluminam riscos importantes: (i) Ao depender da validação da IA, o estudante pode desenvolver uma autoestima acadêmica fragilizada, acreditando estar correto apenas quando a IA confirma; (ii) Tais reflexões devem ser consideradas no momento de levantar os requisitos da solução proposta.

As páginas 3 e 4 são destinadas à formalização dos requisitos da solução tecnológica. Nessa etapa, o protótipo conceitual é traduzido em especificações funcionais, incorporando tanto as implicações levantadas quanto as reflexões produzidas após a exploração do Cyborg AI. Cada requisito apresenta uma descrição objetiva e uma explicação que explicita sua relação com os questionamentos emergidos da interação, garantindo que o *design* da solução considere os elementos discutidos.

No exemplo ilustrativo, é elaborado um conjunto de seis requisitos funcionais. Contudo, não se estabelece que os requisitos devam, necessariamente, pertencer a uma categoria específica. A Figura 1.13 apresenta os três primeiros requisitos registrados no caderno.

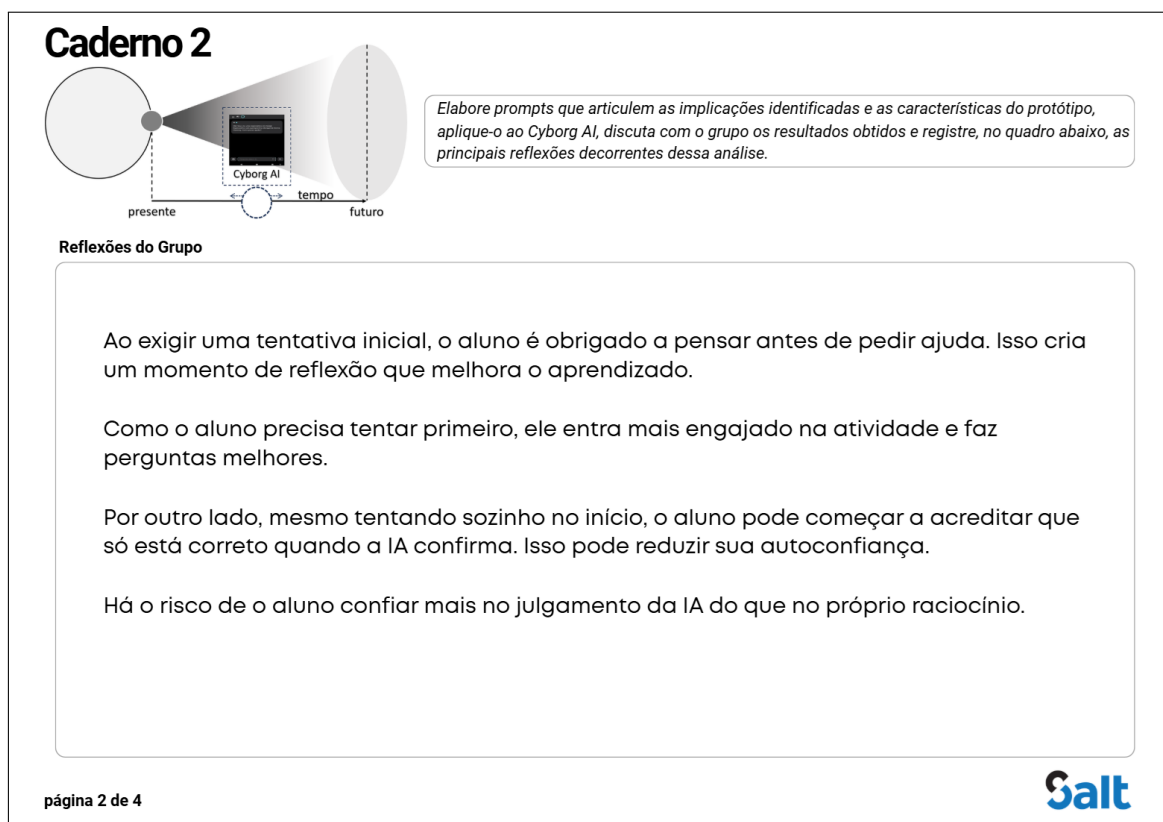


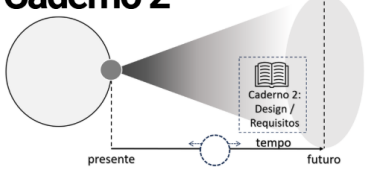
Figura 1.12: Caderno 2 – página 2

Em algumas situações, estabeleceu-se a relação entre os requisitos e as reflexões levantadas na página 2 do mesmo caderno de apoio. Por exemplo, o RF02 está diretamente relacionado à reflexão de que exigir uma tentativa inicial cria um “momento de reflexão” que melhora o aprendizado. Durante a exploração do Cyborg AI, identificou-se que, ao registrar sua primeira ideia, o estudante é levado a organizar pensamentos, arriscar hipóteses e engajar-se cognitivamente antes de receber qualquer apoio algorítmico. Esse espaço de elaboração própria reduz o uso impulsivo da IA e favorece perguntas mais qualificadas. Assim, o RF02 transforma essa reflexão em uma regra técnica que assegura que o aprendizado permaneça ativo e autoral desde o início da interação.

Já o RF03 está relacionado à reflexão que alerta para o risco de o estudante desenvolver dependência cognitiva e confiar mais no julgamento da IA do que no próprio raciocínio. Ao bloquear pedidos diretos de respostas completas e devolver a mensagem “Mostre sua ideia primeiro”, o requisito atua como mecanismo de contenção pedagógica, impedindo o comportamento de consumo passivo da IA identificado nas reflexões com o Cyborg AI. Assim, RF03 protege a autonomia intelectual do aluno e evita que a IA ocupe o papel de autoridade absoluta, reforçando a necessidade de participação ativa e produção própria em todas as etapas do processo. A Figura 1.14 apresenta os requisitos adicionais que compõem o exemplo ilustrativo.

O RF05 está diretamente relacionado à reflexão de que a aprendizagem não deve ser entendida como uma checagem única de certo ou errado, mas como um processo con-

Caderno 2



Elabore os requisitos da solução tecnológica levando em conta as implicações, o protótipo desenvolvido e as reflexões produzidas pelo grupo durante a exploração do Cyborg AI.

Não é necessário classificá-los como funcionais ou não funcionais. Considere que a exploração do Cyborg AI pode ter transformado sua percepção sobre o protótipo. Assim, explique de que maneira cada requisito se conecta às reflexões produzidas durante a exploração do Cyborg AI.

Requisito

RF01 – Registro de atividade e objetivos de aprendizagem
 O sistema deve permitir que professores cadastrem atividades com:
 a) enunciado do problema;
 b) objetivos de aprendizagem;
 c) critérios de avaliação (conceituais, procedimentais, atitudinais).

RF02 – Obrigatoriedade de tentativa inicial
 Antes de acessar qualquer apoio da IA para uma atividade, o estudante deve obrigatoriamente registrar uma tentativa inicial (texto, código, esboço, resposta parcial etc.).
Relação com minhas reflexões: garante que o aluno pense antes de pedir ajuda e cria o “momento de reflexão”.

RF03 – Bloqueio de pedidos diretos de solução
 Se o estudante tentar solicitar diretamente a resposta completa sem registrar tentativa (ou sem atualizar sua tentativa), o sistema deve bloquear o pedido e retornar a mensagem padronizada: “Mostre sua ideia primeiro para que eu possa te ajudar melhor.”
Relação com minhas reflexões: reduz o comportamento de consumo passivo de respostas prontas.

página 3 de 4

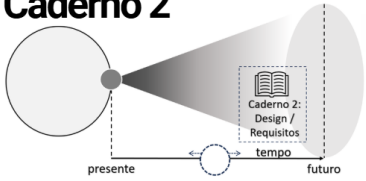

Figura 1.13: Caderno 2 – página 3

tínuo de elaboração. Durante a exploração do Cyborg AI, emergiu a percepção de que o estudante pode se engajar de forma mais profunda quando revisa sua resposta, elabora novas versões e solicita novos feedbacks; esse movimento iterativo fortalece o desenvolvimento do raciocínio e reduz a dependência de uma validação imediata da IA. O RF05 incorpora essa reflexão ao permitir múltiplos ciclos de tentativa e revisão, registrando o histórico de versões e promovendo uma aprendizagem progressiva, situada e autoral.

O RF06 responde diretamente à reflexão que evidenciou o risco de o estudante confiar excessivamente no julgamento da IA, em detrimento de sua própria agência cognitiva. A exploração com o Cyborg AI mostrou que, sem estímulos para questionar, discordar ou justificar posicionamentos, o aluno pode aceitar passivamente as sugestões oferecidas. O RF06 operacionaliza essa preocupação ao introduzir perguntas que provocam o juízo crítico, como: “Você concorda? Por quê?”; incentivando o estudante a avaliar as sugestões da IA e a se posicionar. Com isso, o requisito atua como contrapeso à dependência algorítmica, reforçando a autonomia intelectual e promovendo uma postura reflexiva e argumentativa.

Embora os requisitos ilustrativos estabeleçam relação direta com as reflexões produzidas, essa vinculação não é obrigatória. Os requisitos também podem ser fundamentados nas implicações positivas e negativas identificadas no Caderno 1, ou ainda no próprio *design* da solução. Contudo, as reflexões geradas a partir da interação com o Cyborg AI oferecem um respaldo argumentativo mais robusto, fortalecendo a justificativa e a consis-

Caderno 2



Elabore os requisitos da solução tecnológica levando em conta as implicações, o protótipo desenvolvido e as reflexões produzidas pelo grupo durante a exploração do Cyborg AI.

Não é necessário classificá-los como funcionais ou não funcionais. Considere que a exploração do Cyborg AI pode ter transformado sua percepção sobre o protótipo. Assim, explique de que maneira cada requisito se conecta às reflexões produzidas durante a exploração do Cyborg AI.

Requisito

RF04 – IA responde sempre com base na produção do estudante
 A IA deve gerar feedback somente a partir do conteúdo produzido pelo estudante (tentativa inicial e revisões), evitando fornecer uma solução completa pronta. O feedback deve:

- a) apontar caminhos, não respostas finais;
- b) destacar trechos da resposta do estudante;
- c) sugerir revisões, exemplos, perguntas guias.

RF05 – Múltiplos ciclos de tentativa-feedback
 O sistema deve permitir que o estudante:

- a) revise sua resposta após o feedback;
- b) registre uma nova versão;
- c) peça novo feedback da IA com base nessa versão.

O histórico de versões deve ser armazenado.

Relação com minhas reflexões: reforça o processo iterativo de aprendizagem, e não uma checagem única de certo/errado.

RF06 – Mecanismo de divergência saudável (encorajar juízo próprio)
 A IA deve, em parte das interações, devolver perguntas do tipo:

- a) “Você concorda com essa sugestão? Por quê?”
- b) “Você vê algum problema nessa abordagem que eu propus?”

O sistema deve registrar quando o estudante concorda ou discorda, incentivando que ele se posicione.

Relação com minhas reflexões: combate a confiança cega na IA e reforça a agência do aluno.

página 4 de 4



Figura 1.14: Caderno 2 – página 4

tência dos requisitos elaborados.

1.6.4. Cyborg AI

O Cyborg AI é uma ferramenta desenvolvida para apoiar a reflexão sobre requisitos de software para além das abordagens tradicionais de elicitación e especificação. Seu objetivo é favorecer a construção de requisitos mais socialmente conscientes, por meio de interações baseadas em *prompts* nos quais o usuário descreve características de design consideradas em sua proposta. A partir dessas informações, o Cyborg AI gera respostas orientadas por um conjunto de temas centrais do pensamento de Donna Haraway.

Entre esses temas destacam-se a problematização do que é tomado como natural, a compreensão dos hibridismos como condição constitutiva da contemporaneidade, a necessidade de considerar relações multiespécies e de coabitação, a valorização dos saberes situados e a figura do ciborgue como instrumento de libertação e imaginação política na projeção de futuros desejáveis. A ferramenta também incorpora uma perspectiva crítica sobre a construção material do poder, evidenciando como decisões técnicas participam da produção de assimetrias sociotécnicas.

Dessa forma, o Cyborg AI atua como um dispositivo reflexivo que amplia o espaço de problematização ética, política e social no processo de definição de requisitos de software.

1.7. Considerações Finais

Neste minicurso argumentamos que o paradigma da bifurcação, herdado da modernidade, limita a capacidade de compreender e intervir criticamente em ecossistemas sociotécnicos cada vez mais complexos, híbridos e opacos. Neste sentido, propomos uma reorientação conceitual e prática sobre como projetamos, analisamos e avaliamos tecnologias.

O contato direto com as ferramentas do *Speculative Entangled Design* (SpED) / Sociotechnical Entanglement Framework (SEF) possibilita vivenciar, de forma coletiva e situada, a mudança de uma lógica instrumental para uma lógica relacional de *design*. Essa vivência permite perceber, na prática, como decisões técnicas materializam valores e produzem efeitos éticos e políticos.

O acesso a este material por quem não pôde participar do minicurso é relevante. O texto organiza os fundamentos teóricos, metodológicos e conceituais do SpED/SEF combinados com temas explorados por Donna Haraway, oferecendo um percurso claro para compreender por que abordagens tradicionais de design e engenharia de software são insuficientes frente à IA e a sistemas sociotécnicos avançados. Ao apresentar os princípios que fundamentam o SpED/SEF, as ferramentas e exemplos de aplicação, o material atua como um guia reflexivo que permite ao leitor realizar, de forma autônoma, parte do exercício crítico desenvolvido no minicurso.

O participante do minicurso (ou leitor deste material) é convidado a questionar a neutralidade dos sistemas, a reconhecer a agência dos não humanos, a lidar com a incerteza por meio da especulação e a assumir a responsabilidade inerente a toda prática de *design*. Em um cenário marcado pela crescente automação de decisões, pela opacidade algorítmica e pela intensificação das assimetrias sociotécnicas, essa reflexão deixa de ser opcional e passa a ser uma condição fundamental para a atuação responsável em SI.

Agradecimentos

Este estudo foi parcialmente financiado pelo Conselho Nacional de Desenvolvimento Científico e Tecnológico (CNPq), por meio do Processo nº 313783/2025-0, pela Fundação Carlos Chagas Filho de Amparo à Pesquisa do Estado do Rio de Janeiro (FAPERJ), Processos nº E-26/210.792/2024 e nº 316006, e pela Coordenação de Aperfeiçoamento de Pessoal de Nível Superior (CAPES), no âmbito do Programa de Apoio à Disseminação de Informação Científica e Tecnológica (PADICT) e do Portal de Periódicos da CAPES. O aluno Gabriel participou desta pesquisa como bolsista de Iniciação Científica do Programa Institucional de Bolsas de Iniciação Científica (PIBIC).

Referências

- Bardzell, J. and Bardzell, S. (2013). What is "critical" about critical design? In *Proceedings of the SIGCHI conference on human factors in computing systems*, pages 3297–3306.
- Candy, S. and Dunagan, J. (2017). Designing an experiential scenario: The people who vanished. *Futures*, 86:136–153.
- Chamaidi, T. and Stavrakis, M. (2024). A multispecies interaction design approach:

- Introducing the beings activities context technologies (bact) framework. *Multimodal Technologies and Interaction*, 8(9):77.
- Chisita, C. T., Durodolu, O. O., and Rusero, A. M. (2025). Data capitalism in the milieu of the surveillance economy: What can libraries do? *IFLA Journal*, 51(2):339–349.
- da Silva Paulichi, J. (2025). The digital persona and surveillance capitalism: contemporary challenges to the protection of personality rights in cyberspace1. *Pensar*, 30:1–20.
- Dunne, A. and Raby, F. (2024). *Speculative Everything, With a new preface by the authors: Design, Fiction, and Social Dreaming*. MIT press.
- Fitzgerald, A. (2025). Death by data: Abstraction and the political economy of computationally driven state violence. *Media Theory*, 9(1):169–200.
- Forlano, L. E. and Halpern, M. K. (2023). Speculative histories, just futures: From counterfactual artifacts to counterfactual actions. *ACM Transactions on Computer-Human Interaction*, 30(2):1–37.
- Frauenberger, C. (2019). Entanglement hci the next wave? *ACM Transactions on Computer-Human Interaction (TOCHI)*, 27(1):1–27.
- Giaccardi, E., Redström, J., and Nicenboim, I. (2025). The making (s) of more-than-human design: introduction to the special issue on more-than-human design and hci. *Human-Computer Interaction*, 40(1-4):1–16.
- Haraway, D. (2013). A cyborg manifesto: Science, technology, and socialist-feminism in the late twentieth century. In *The transgender studies reader*, pages 103–118. Routledge.
- Harman, G. (2018). *Object-oriented ontology: A new theory of everything*. Penguin UK.
- Hines, A., Bishop, P. J., and Slaughter, R. A. (2006). *Thinking about the future: Guidelines for strategic foresight*. Social Technologies Washington, DC.
- Latour, B. (2005). *Reassembling the social: An introduction to actor-network-theory*. Oxford university press.
- Latour, B. (2012). *We have never been modern*. Harvard university press.
- Loutfi, M. S. (2025). *Entanglement Theories and Speculative Design in the Co-Constitution of Sociotechnical Realities*. Tese de doutorado, Universidade Federal do Estado do Rio de Janeiro, Rio de Janeiro, RJ, Brasil.
- Luo, H. and Shi, D. (2025). Algorithmic embodiment and the posthuman condition: Reframing human autonomy in n. katherine hayles’ media culture theory. *Convergence*, page 13548565251403045.
- Lupton, D. (2015). Donna haraway: The digital cyborg assemblage and the new digital health technologies. In *The Palgrave handbook of social theory in health, illness and medicine*, pages 567–581. Springer.

- Malpass, M. (2019). *Critical design in context: History, theory, and practice*. Bloomsbury Publishing.
- Milojević, I. and Inayatullah, S. (2015). Narrative foresight. *Futures*, 73:151–162.
- Mitrović, I., Auger, J., Hanna, J., and Helgason, I. (2021). *Beyond speculative design: past–present–future*. SpeculativeEdu; Arts Academy, University of Split.
- Østbø, J. (2021). Hybrid surveillance capitalism: Sber’s model for russia’s modernization. *Post-Soviet Affairs*, 37(5):435–452.
- Pasa, B. and Sinni, G. (2024). Democracy in outer space: Speculative design for future citizenship. In *For Nature/With Nature: New Sustainable Design Scenarios*, pages 965–980. Springer.
- Petronio, R. (2023). Introdução à teoria gerativa—parte 1:: Conhecimento, cosmologia e emergência a partir da obra de david deutsch. *TECCOGS: Revista Digital de Tecnologias Cognitivas*, 1(27).
- Qi, Y., Sun, Y., Liu, Q., Zhang, Q., Cai, H., and Zheng, Q. (2023). The intersection of artificial intelligence and brain for high-performance neuroprosthesis and cyborg systems. *Frontiers in Neuroscience*, 17:1133002.
- Rafael, S., Silva, B., Anjos, H., Meintjes, L., and Tavares, P. (2023). Data surveillance in capitalism society: The globule app, a speculative design to control the algorithm. In *Proceedings of the 2023 ACM International Conference on Interactive Media Experiences Workshops*, pages 27–31.
- Solhjoo, N. (2024). Knowing within multispecies families: An information experience study. *Journal of Information Science*, page 01655515241268845.
- Stojanović, Đ. (2023). From cyborg to cybernantrope: Basic political, cultural and philosophical dimensions of the concepts. *Politika nacionalne bezbednosti*, 25(2):139–153.
- Suchman, L. (2002). Located accountabilities in technology production. *Scandinavian journal of information systems*, 14(2):7.
- van der Linden, D. (2021). Interspecies information systems. *Requirements Engineering*, 26(4):535–556.
- Verbeek, P.-P. (2012). Expanding mediation theory. *Foundations of science*, 17(4):391–395.
- Whitehead, A. N. and Douchement, J. (1920). *The concept of nature*, volume 190. Springer.
- Winner, L. (1986). *Computer Ethics*, chapter Do artifacts have politics? The University of Chicago Press.

Zuboff, S. (2022). Surveillance capitalism or democracy? the death match of institutional orders and the politics of knowledge in our information civilization. *Organization Theory*, 3(3):26317877221129290.

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 ~~future~~ 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.

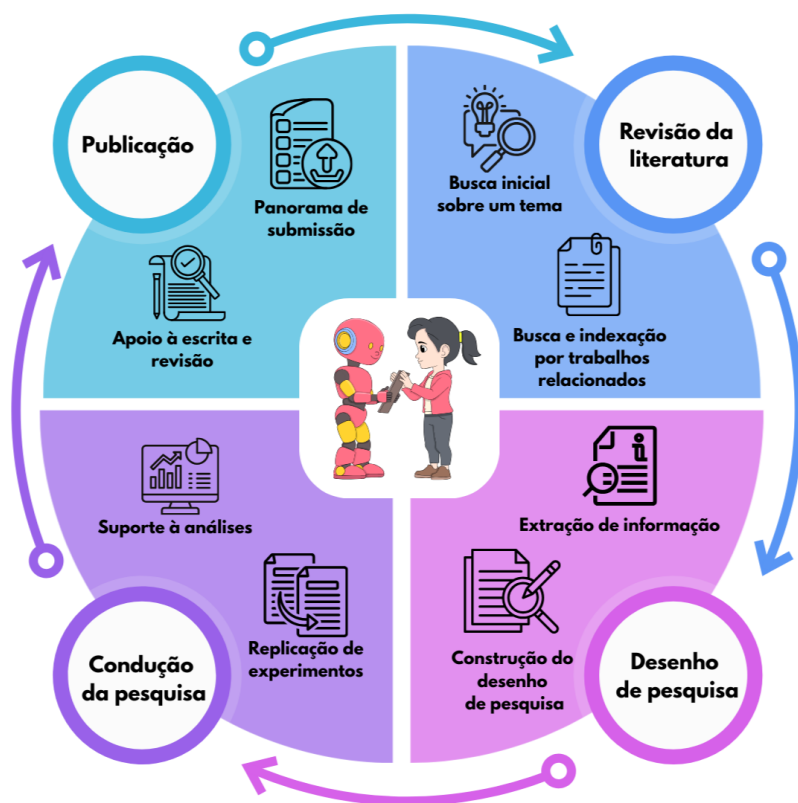


Figura 2.1: Taxonomia de tarefas de pesquisa utilizando a IA como um co-cientista, proposta neste trabalho.

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 ChatGPT¹ 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.

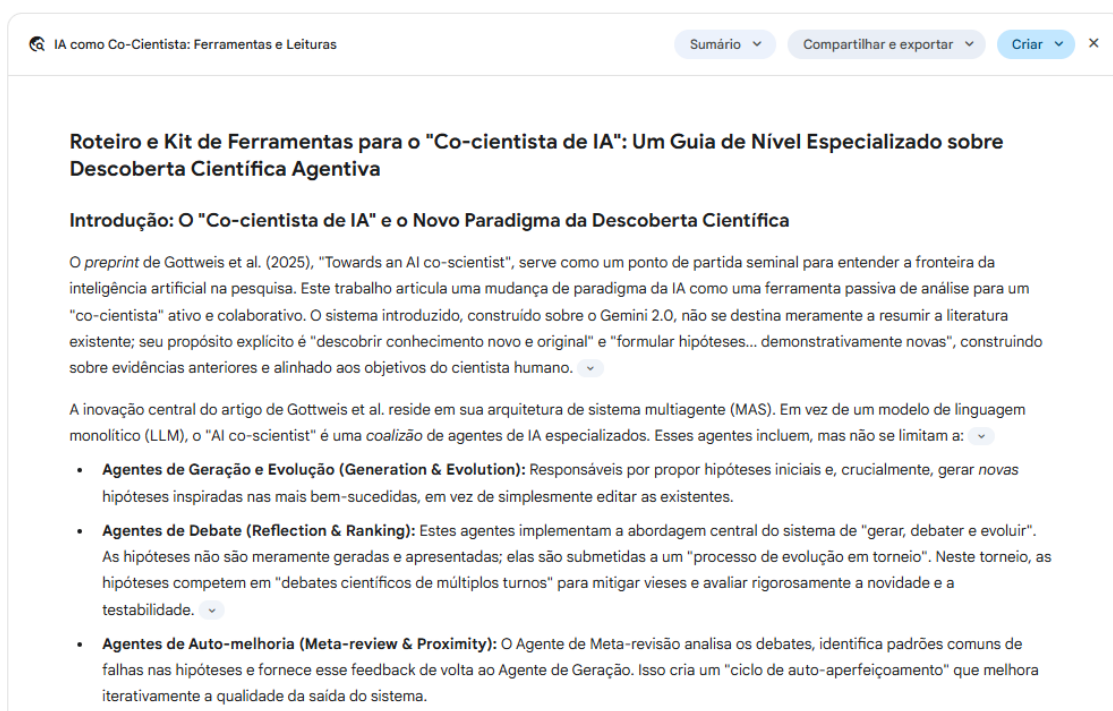


Figura 2.4: Relatório gerado com o *Deep Research* do Gemini.

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.

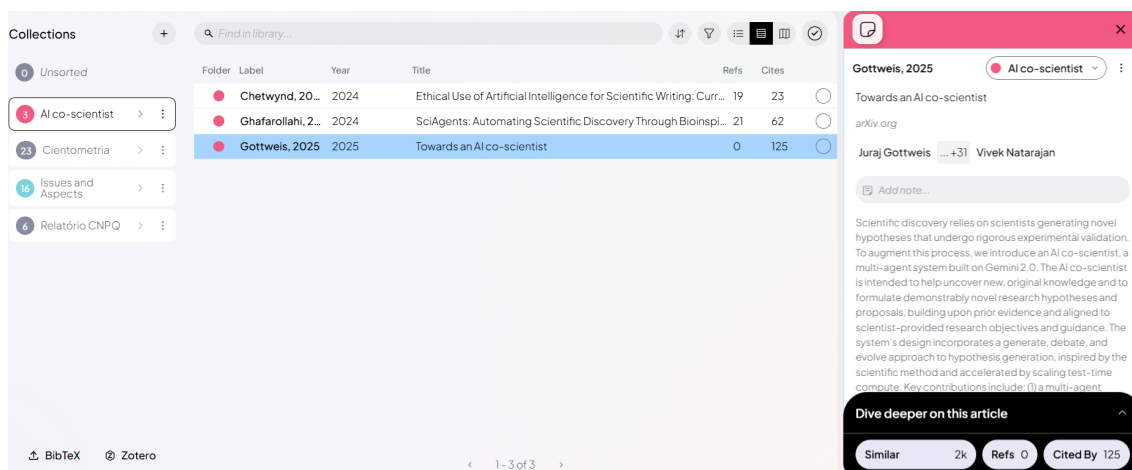


Figura 2.5: Coleção de artigos organizada no ResearchRabbit.

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

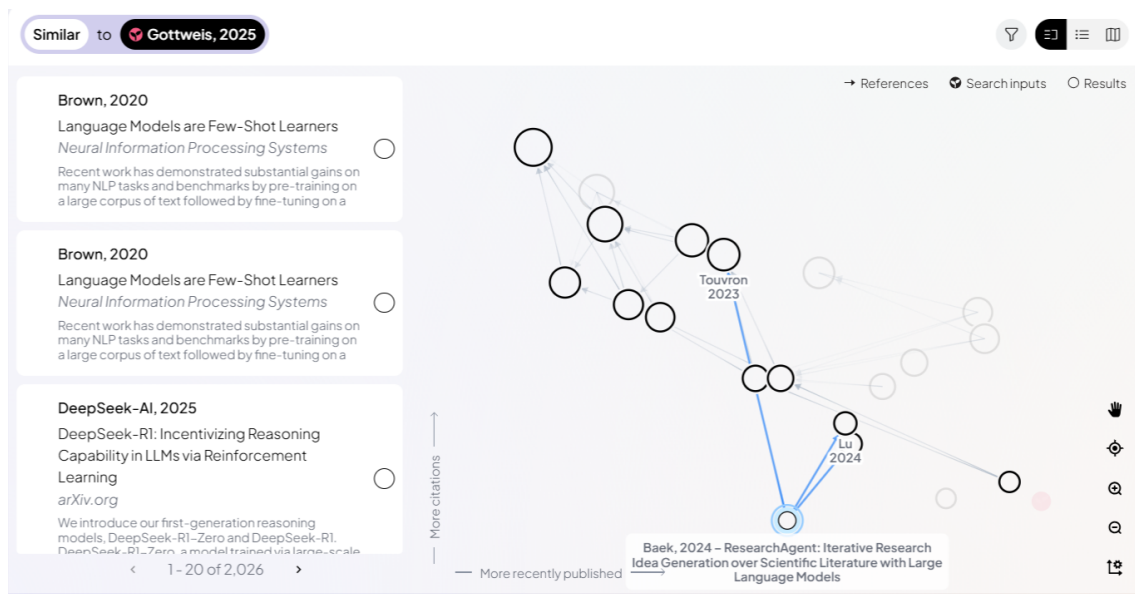


Figura 2.6: Artigos similares indicados pelo ResearchRabbit.

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.

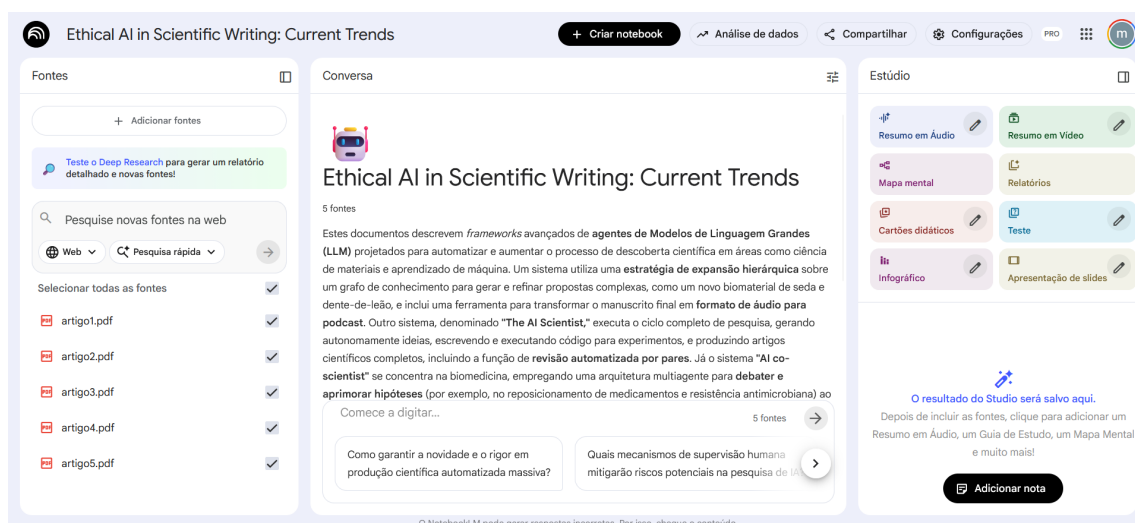


Figura 2.7: Tela inicial do NotebookLM.

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.

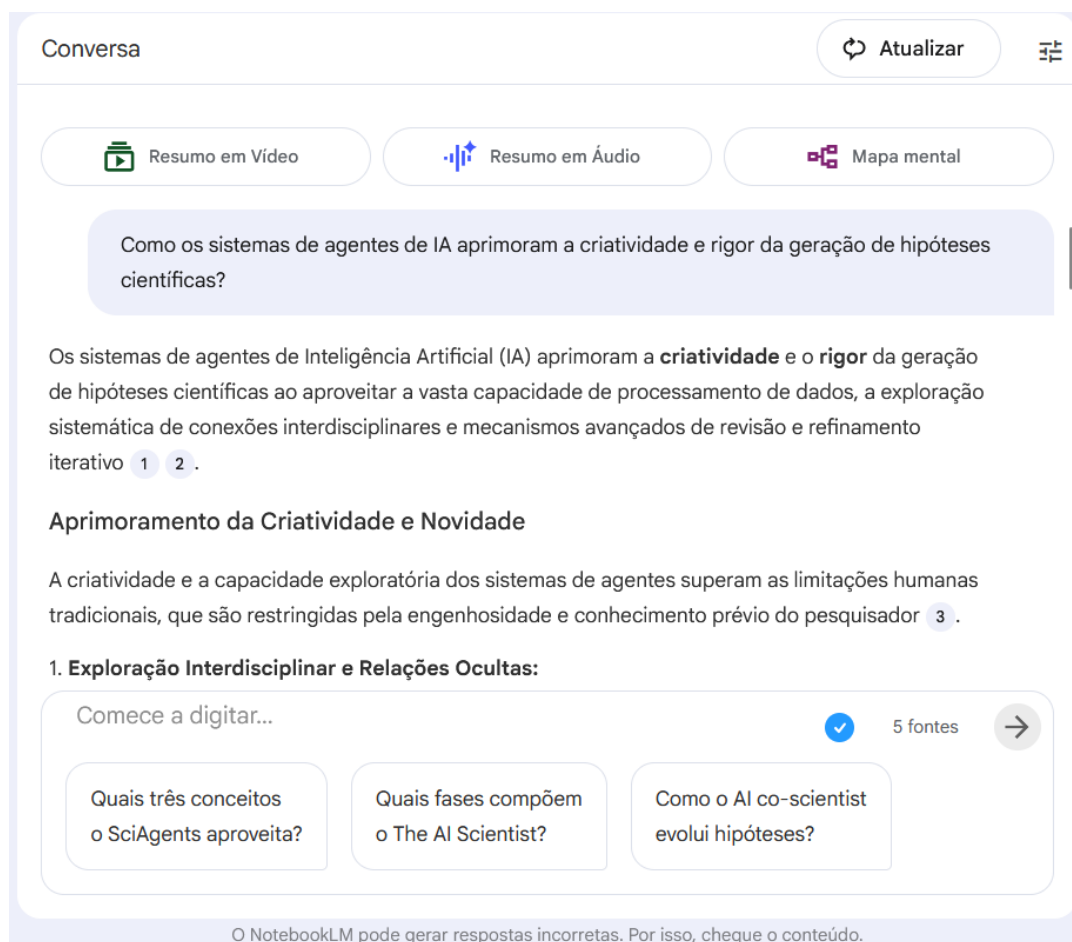


Figura 2.8: Resposta a uma requisição no NotebookLM.

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.

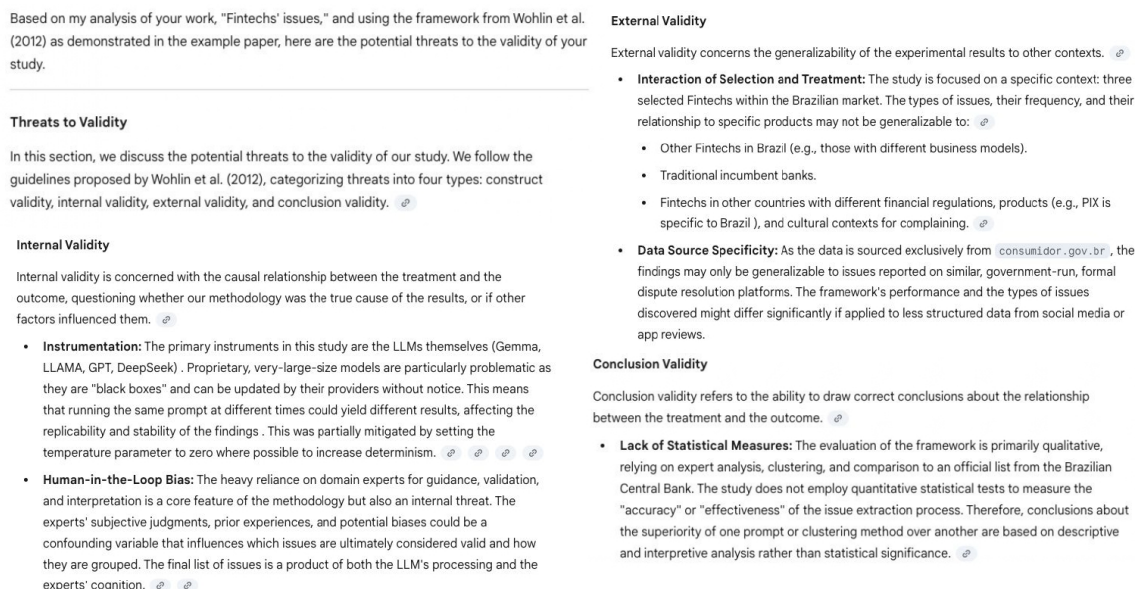


Figura 2.14: Resposta do Gemini referente ao *prompt* para identificar ameaças à validade da pesquisa.

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

Com base no artigo científico "BraSNAM em perspectiva: uma análise da sua trajetória até os 10 anos de existência", foram extraídas as seguintes etapas técnicas implementáveis:

Passos principais (macroetapas)

1. **Coletar** os dados básicos dos artigos (título, palavras-chave, autores, instituições e ano de publicação) da base SBC OpenLib (SOL) utilizando um *web crawler*.
2. **Coletar** manualmente o número de citações de cada artigo a partir do título na plataforma Google Acadêmico.
3. **Consultar** a API "Nomes" do IBGE para inferir o gênero dos autores a partir do primeiro nome.
4. **Pré-processar** os dados textuais (palavras-chave, autores e instituições) para padronização.
5. **Padronizar** as palavras-chave com base na similaridade de strings Levenshtein (superior a 85%), com revisão manual.
6. **Padronizar** os nomes dos autores, mantendo apenas o primeiro e o último nome e ajustando abreviações (ex: "Jr." para "Júnior").
7. **Padronizar** os nomes das instituições por meio de mapeamento manual.
8. **Gerar** listas de adjacências para coocorrência de palavras-chave, coautoria e colaborações interinstitucionais, pareando os dados por artigo.
9. **Visualizar** e analisar os componentes das redes geradas (coautoria e colaborações interinstitucionais).

Para o Passo 1:

- "Crie um script em Python utilizando a biblioteca BeautifulSoup para fazer a varredura do arquivo de anais do BraSNAM na SBC OpenLib, disponível em <https://sol.sbc.org.br/index.php/brosonam/issue/archives>. O script deve extrair as seguintes informações de cada artigo: título, palavras-chave, lista de autores, lista de instituições e ano de publicação. Salve o resultado em um arquivo CSV."

Para o Passo 2:

- "Desenvolva uma função em Python que receba uma lista de títulos de artigos e, para cada título, construa uma URL de busca para o Google Acadêmico e a exiba, instruindo o usuário a inserir manualmente o número de citações encontrado."

Para o Passo 3:

- "Escreva um script Python que leia uma lista de nomes de autores de um arquivo CSV, extraia o primeiro nome de cada um e faça uma requisição à API 'Nomes' do IBGE (<https://servicodados.ibge.gov.br/api/v2/censos/nomes/>) para determinar o gênero. Adicione uma nova coluna ao CSV com o gênero inferido ou 'indefinido' caso o nome não seja encontrado."

Para o Passo 4:

- "Gere uma função em Python que receba uma string como entrada e execute as seguintes operações de limpeza: remoção de espaços em branco duplicados, remoção de espaços no início e no fim da string, remoção de caracteres especiais e conversão de todo o texto para caixa-baixa."

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.

Com certeza. O prompt para o Passo 6 focará na replicação exata das duas regras de padronização de nomes de autores mencionadas no artigo.

Prompt para Geração de Código - Passo 6

Persona: Você é um assistente de pesquisa focado na replicação de metodologias científicas, responsável por criar um script em Python para a etapa de padronização de nomes de autores.

Tarefa: Crie um script em Python para padronizar os nomes dos autores de um conjunto de dados, implementando estritamente as regras de transformação descritas no artigo "BraSNAM em perspectiva...".

Fonte da Metodologia: O script deve ser uma implementação direta dos procedimentos descritos na seção 2.3 do artigo. A fonte especifica duas operações: "Por exemplo, Jr. foi padronizado para Júnior" e "Além disso, consideramos apenas o primeiro e o último nome".

Requisitos Funcionais:

1. **Linguagem:** O script deve ser escrito em Python.
2. **Arquivo de Entrada:**
 - O script deve ler o arquivo `dados_brasnom_preprocessados.csv` (gerado no Passo 4), que já contém os nomes em caixa-baixa e com limpeza básica.
3. **Coluna-Alvo:**
 - O processo de padronização deve ser aplicado à coluna `autores`.

4. Lógica de Padronização:

- O script deve processar a coluna `autores`, que contém um ou mais nomes separados por ponto e vírgula (;).
- Para cada célula da coluna `autores`: a. Separe a string em uma lista de nomes individuais. b. Para cada nome individual na lista, aplique as seguintes regras na ordem: i. **Substituição de Termos:** Encontre e substitua a sub-string "jr." (em caixa-baixa) por "júnior". ii. **Redução do Nome:** Divida o nome em componentes (palavras) e construa uma nova string contendo apenas o primeiro e o último componente, separados por um espaço. Por exemplo, "fábio m. f. lobato" deve se tornar "fábio lobato". c. Junte a lista de nomes processados de volta em uma única string, usando o ponto e vírgula (;) como separador.
- Substitua o valor original da célula pelo novo valor padronizado.

Formato da Saída:

1. **Arquivo de Saída:** Salve o DataFrame com os nomes dos autores padronizados em um novo arquivo CSV. Um nome apropriado seria `dados_brasnom_autores_padronizados.csv`.
2. **Estrutura:** O arquivo de saída deve manter a mesma estrutura do arquivo de entrada, mas com os valores da coluna `autores` atualizados conforme a metodologia.

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 Selecionados:** 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 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 GitHub¹⁰.

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/fabiolobato/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.