

IA GENERATIVA NA ANÁLISE DE DADOS QUALITATIVOS: um olhar crítico para a pesquisa em linguística aplicada

GENERATIVE AI IN QUALITATIVE DATA ANALYSIS: a critical approach for research in applied linguistics

Flávia Bonella Ribeiro¹

Cláudia Jotto Kawachi-Furlan²

Antônio A. de A. Rocha³

RESUMO: Em face de uma revolução tecnológica sem precedentes, percebe-se que as IAs generativas estão alterando substancialmente as atividades pessoais e profissionais, e o fazer científico não é exceção. O uso das IAs generativas para a pesquisa impõe uma série de questões que trazem implicações quanto à ética, validade e confiabilidade dos trabalhos. Entretanto, existe uma escassez de registros e análises sistematizadas do uso dessas ferramentas, em especial, na pesquisa qualitativa. Dessa forma, o presente artigo tem como objetivo descrever a aplicação de tais ferramentas em uma etapa crucial da pesquisa qualitativa – a análise de dados textuais – e problematizar seu uso no campo da Linguística Aplicada. Para investigar esta questão, foram empregadas três soluções de IA generativa – ChatGPT, ATLAS.ti e Manus – na análise de um subconjunto de uma coleção de dados mais ampla. Os resultados obtidos demonstram o potencial dessas ferramentas quando somado à supervisão humana, com destaque para a plataforma Manus, para a análise qualitativa dos dados. Simultaneamente, os achados reforçam a imperatividade de uma análise crítica e criteriosa das respostas geradas pelos modelos de IA generativa.

Palavras-chave: Pesquisa qualitativa. Linguística Aplicada. Inteligência Artificial.

ABSTRACT: In the face of an unprecedented technological revolution, it is evident that generative AIs are substantially transforming personal and professional activities, and scientific inquiry is no exception. The use of generative AI in research raises a series of issues with implications for research ethics, validity, and reliability. However, there is a lack of systematic entries and analyses regarding the use of these tools, particularly in qualitative research. Thus, the aim of this article is to document and analyze the application of such tools in a crucial stage of qualitative research – the analysis of textual data – and to problematize their use within the field of Applied Linguistics. To address this question, three generative AI solutions – ChatGPT, ATLAS.ti, and Manus – were employed in the analysis of a subset of a broader data collection. The results demonstrate the potential of these tools when combined with human supervision, with particular emphasis on the Manus platform, for qualitative data analysis. At the same time, the findings reinforce the need of conducting a critical and rigorous examination of the responses generated by generative AI models.

Keywords: Qualitative research. Applied Linguistics. Artificial Intelligence.

¹ Instituto Federal do Espírito Santo, Cariacica, Espírito Santo, Brasil. flavia.bonella@ifes.edu.br.

² Universidade Federal do Espírito Santo, Departamento de Línguas e Letras, Programa de Pós-Graduação em Linguística da Universidade Federal do Espírito Santo, Vitória, Espírito Santo, Brasil. claudia.furlan@ufes.br.

³ Instituto de Computação da Universidade Federal Fluminense, Niterói, Rio de Janeiro, Brasil. arocha@ic.uff.br.

INTRODUÇÃO

A aplicação de Inteligência Artificial (IA) generativa em processos educacionais e na pesquisa em educação tem se expandido exponencialmente. Em face de uma revolução tecnológica sem precedentes, Silva (2024) reitera a necessidade de uma revisão das práticas vigentes e de uma busca pela compreensão deste novo paradigma. Nesse contexto, pesquisadores no campo da Linguística Aplicada têm direcionado seu foco para os impactos do uso da IA nos processos de ensino e aprendizagem (D'Esposito; Gatner, 2024; Martins; Pitz, 2024). Contudo, observa-se uma lacuna na literatura no que tange à investigação dos impactos da IA generativa especificamente no âmbito da análise de dados de pesquisas nessa área.

Embora relatos informais sobre a experimentação e o uso de ferramentas de IA por pesquisadores e pós-graduandos sejam prevalentes, há uma escassez de registros e análises sistematizadas desses resultados. Considerando que o uso das IAs generativas para o trabalho científico impõe uma série de questões que trazem implicações quanto à ética, a validade e a confiabilidade da pesquisa, o presente artigo tem como objetivo descrever a aplicação de tais ferramentas em uma etapa crucial da pesquisa qualitativa – a análise de dados textuais – e problematizar seu uso no campo da Linguística Aplicada.

Nesse sentido, este trabalho busca discutir de que forma a IA generativa pode contribuir para a análise de dados qualitativos em pesquisas na área de Linguística Aplicada. Indagamos, portanto, se as ferramentas de IA generativa e os modelos de linguagem estão aptos a identificar, em uma transcrição de entrevista com uma docente, práticas que possam ser associadas às teorias de Letramento Crítico (LC) e justificar de forma coerente o embasamento teórico que permitiu tal associação.

A estrutura deste artigo organiza-se da seguinte forma: a Seção 2 oferece uma revisão teórica sobre IA generativa e sua aplicação em pesquisa qualitativa. A Seção 3 descreve a metodologia dos experimentos, detalhando o uso das três ferramentas de IA e os dados utilizados. A Seção 4 apresenta uma análise crítica dos resultados da análise de dados gerada pelas plataformas. Por fim, a Seção 5 expõe as considerações finais do estudo.

A IA GENERATIVA E A PESQUISA QUALITATIVA: UMA REVISÃO TEÓRICA

A Inteligência Artificial tem evoluído para além de suas aplicações tradicionais, que se concentravam em tarefas específicas como a identificação de padrões e a realização de previsões com

base em dados históricos. Uma nova fronteira, a IA generativa, emergiu com a capacidade de criar conteúdo original em resposta a comandos do usuário, conhecidos como prompts. Esta subárea da IA foca no desenvolvimento de modelos capazes de gerar conteúdo sem a necessidade de programação ou instruções explícitas. A criação de conteúdo é derivada de padrões aprendidos a partir de vastos volumes de dados. Para tal, as IAs generativas empregam técnicas avançadas de aprendizado de máquina, como o aprendizado profundo (*deep learning*) e as redes neurais, para construir modelos que podem produzir uma variedade de mídias, incluindo textos, imagens, sons e vídeos. Estas são apenas algumas das inúmeras aplicações possíveis para as IAs generativas, que estão redefinindo os limites da criação de conteúdo digital.

Dentre os diversos modelos de IA generativa, os Grandes Modelos de Linguagem (LLMs, do inglês Large Language Models) se destacam. Estes modelos utilizam arquiteturas de redes neurais profundas, mais especificamente, arquiteturas baseadas em Transformers (Vaswani *et al.*, 2017), para processar e gerar linguagem natural de forma contextual, coerente e relevante. A arquitetura Transformer revolucionou o processamento de linguagem natural ao dispensar o uso de recorrência e convoluções, baseando-se unicamente em mecanismos de atenção para capturar dependências de longo alcance nos dados. Os LLMs são a base para a maioria das IAs generativas preeminentes, capacitando-as a gerar uma ampla gama de conteúdos. Ferramentas como o Gemini da Google, o CoPilot da Microsoft, o ChatGPT da OpenAI e Manus da Butterfly Effect Technology são exemplos notáveis que surgiram no mercado nos últimos anos, demonstrando o poder e a versatilidade desses modelos.

Embora a IA generativa já seja uma realidade no cotidiano e esteja sendo aplicada em diversas atividades, é crucial reconhecer que tanto as ferramentas quanto seus modelos subjacentes ainda estão em fase de desenvolvimento e evolução tecnológica. Consequentemente, existem questões técnicas, éticas e legais que demandam consideração cuidadosa em seu uso. Uma dessas limitações técnicas bem notória das ferramentas de IA generativa é conhecida como o fenômeno das alucinações e decorre do fato de que esses modelos são desenvolvidos para produzir algum resultado de geração de conteúdo a partir de comandos de entrada, por isso a ferramenta pode gerar conteúdos imprecisos ou mesmo incorretos, ainda que pareçam coerentes.

Além disso, há de considerar que as empresas responsáveis pela criação dessas tecnologias têm relação direta com preceitos neoliberais e buscam, como fim, o lucro. Assim, destacam-se os vieses empregados nas programações e algoritmos visando à produção de resultados tendenciosos, ou seja, que servem ao capital, bem como às estruturas de poder, especialmente ao poder político. Ao tratarmos do uso da IA generativa como ferramenta para análise de dados qualitativos, estamos cientes dessas questões e, dessa forma, agindo com a realidade que está posta e com as implicações, nunca neutras,

das concepções de língua e linguagem relacionadas com este trabalho e dos usos de tecnologias no momento atual.

Portanto, apesar do potencial da IA generativa para revolucionar a forma como criamos e produzimos conteúdo e como podemos otimizar⁴ certas atividades com seu suporte, é imperativo conhecermos as suas reais capacidades e limitações, diante de diferentes contextos, para maximizar seus benefícios e mitigar possíveis riscos.

A pesquisa qualitativa é um terreno complexo. Existe uma gama variada de metodologias que podem ser utilizadas e caminhos que o pesquisador pode percorrer ao realizar uma pesquisa dessa natureza (Sinha *et al.*, 2024). Embora alguns estudos qualitativos sejam centrados principalmente na análise de texto (Prescott *et al.*, 2024), neste trabalho, consideramos as pesquisas que apresentam uma abordagem reflexiva, isto é, pesquisas que envolvem análises profundas e abrangentes e que são realizadas por meio da permanência prolongada em campo, da geração e imersão de uma quantidade substancial de dados e, conseqüentemente, da interpretação desses dados pelo pesquisador (Paulus; Marone, 2025).

Sampaio *et al.* (2024a) discutem amplamente sobre as mudanças trazidas pelas IAs generativas para a pesquisa acadêmica nas áreas das ciências humanas e sociais. Embora a escrita tenha sido o principal foco das discussões até então, os autores trazem para o debate a utilização de IAs para a busca e seleção de literatura acadêmica, leitura de material, análise de dados e tradução. No que tange à análise de dados em pesquisa qualitativa, o texto não fornece exemplos; os autores se limitam a citar o ATLAS.ti e as possíveis potencialidades de interação com os dados. Entretanto, uma importante consideração é feita em relação à qualidade da pesquisa:

O acaso e o oportuno, elementos comuns a empreitadas intelectuais, parecem ficar de lado quando terceirizamos as 'análises' à máquina. Seria possível indagarmos: trata-se de 'análise' ou apenas da correlação de dados em um padrão identificável pela máquina? A resposta a esta questão revela alguns limites do uso científico das IAs: elas são incapazes de reconhecer o que está fora do pensamento computacional (o que equivale a dizer que elas só são capazes de conhecer o que pode ser reduzido ao computacional) e, mais ainda, revelam como as habilidades analíticas humanas têm sido relegadas à mera instrumentalidade, facilmente substituída pelas máquinas (Sampaio *et al.*, 2024a, p. 11).

Embora nosso foco seja a Linguística Aplicada, nos encontramos em consonância com os autores ao entender que as ferramentas de IA ainda não são capazes de reproduzir as subjetividades necessárias para analisar dados qualitativos em sua complexidade. Além disso, os autores alertam para um outro

⁴ Novamente, ressaltamos que essas atividades estão sempre relacionadas com a origem das programações e com as concepções de mundo (e de capital) das grandes empresas de tecnologia. Afinal, otimizar para quem?

problema: a falta de políticas e diretrizes⁵ estabelecidas pelos órgãos que regulam e fomentam a pesquisa no Brasil para garantir a qualidade das pesquisas acadêmicas mediante a popularização das IAs.

Paulus e Marone (2025) investigaram como as maiores empresas, incluindo a ATLAS.ti, descrevem o uso de IA generativa para a pesquisa qualitativa em seus websites e concluem que o discurso usado é incompatível com os fundamentos epistemológicos da pesquisa qualitativa. Esses discursos transformam os pontos fortes desse tipo de pesquisa em problemas a serem mitigados pela IA. Segundo os autores, as empresas oferecem a geração de insights automatizados enquanto o que se busca é uma produção de significação sistemática. Para tanto, é preciso que ocorra um envolvimento grande com os dados, uma atividade que requer um investimento significativo de tempo. As ferramentas, no entanto, oferecem rapidez viabilizada por meio de uma “conversa” com os documentos e não uma análise apropriada dos dados.

O uso de IA favorece o distanciamento entre o pesquisador e os dados. Renunciar à agência humana em nome da rapidez está em oposição à tradição da pesquisa qualitativa que tem o pesquisador como um dos instrumentos de pesquisa (Paulus; Marone, 2025; Morgan, 2023; Bogdan; Biklen, 1994). É ele que tem a vivência junto a comunidade estudada, que participou das atividades diárias, e estabeleceu relações com aqueles indivíduos. Portanto, ao olhar para os dados, é o pesquisador, com base nessa experiência, que estabelece os padrões relevantes, conecta dados de diferentes fontes e assim, chega aos resultados de pesquisa. A análise de dados consiste em um processo interpretativo, cuidadoso e gradual de construção de sentidos por parte do pesquisador.

Pack e Maloney (2023) abordam a questão do uso das IAs generativas na pesquisa em ensino de línguas, demonstrando e discutindo exemplos de como o ChatGPT⁶ da OpenAI pode ser aproveitado como uma ferramenta por pesquisadores. Os autores observaram a forma como a ferramenta executava funções que incluíam a compilação de informação, o fornecimento de citações sobre um tópico específico e o auxílio nas pesquisas por meio de atribuição de tarefas à IA.

Com base nos dados gerados pela interação com as ferramentas, Pack e Maloney (2023) sugerem que as IAs podem ser úteis para a busca e apresentação de informações, assim como podem realizar, de maneira razoável, análises discursivas e linguísticas de textos. Além disso, os pesquisadores disponibilizam interações para mostrar como a IA fornece comentários que muito se assemelham aos comentários de revisores reais e podem auxiliar na construção de uma pesquisa. Contudo, os autores

⁵ Atualmente, as universidades estão em processo de elaboração de resoluções acerca do uso de IA em trabalhos acadêmicos. A CAPES e o CNPq também lançaram guias sobre essa temática, como a Portaria 2664/2026 de 6 de março de 2026 – Política de Integridade na Atividade Científica do CNPq.

⁶ No artigo, os exemplos trazidos pelos autores foram gerados pelo ChatGPT (3.5 Turbo) por ser bastante conhecido e disponibilizar acesso gratuito, mesmo não sendo a versão mais recente.

registram a ocorrência de alucinações ao constatar que citações não existentes foram sugeridas. Outro problema apontado é que o upload de grandes quantidades de dados textuais é uma tarefa trabalhosa, pois a versão testada não era capaz de realizar o comando automaticamente.

Baseados nos pilares da ética em pesquisa na Linguística Aplicada, Pack e Maloney (2023) elencam quatro preceitos com o objetivo de orientar práticas quanto ao uso da IA generativa em pesquisa. Primeiramente, eles afirmam que o processo de pesquisa precisa ser transparente. Isso significa que o uso de ferramentas em qualquer parte do processo de pesquisa precisa ser informado e documentado. O segundo ponto está relacionado à confidencialidade dos participantes. Como empresas terceirizadas terão acesso aos dados gerados, os participantes devem ser informados para que o devido consentimento seja obtido. Em seguida, os autores sugerem que os pesquisadores devem verificar como a área lida com a questão de autoria. A OpenAI não reivindica direitos autorais sobre textos criados pelo GPT, mas eles foram construídos com base em textos já existentes, criados por outros autores. Por fim, os pesquisadores alertam para os vieses e erros presentes nas fontes que treinam as IAs que tendem a ser perpetuados.

Alguns estudos publicados recentemente comparam os resultados de análises feitas por humanos com os resultados obtidos por meio de interações com a IA sobre um mesmo conjunto de dados qualitativos (Morgan, 2023; Hamilton *et al.*, 2023; Chubb, 2023; Prescott *et al.*, 2024). Neste artigo, diferentemente do que observamos nos estudos supracitados, nos propomos a testar a aplicação da IA generativa na análise de dados qualitativos em pesquisa na área da Linguística Aplicada, sem comparar esse conjunto com uma análise realizada por humanos, mas, obviamente sob investigação dos autores deste texto. Entendemos que, neste momento, faz-se necessário explorar e registrar o comportamento das IAs na interação com os dados para que as análises e os resultados sejam validados e assim, práticas desejáveis possam ser identificadas e recomendadas.

METODOLOGIA EXPERIMENTAL REALIZADA

Para a condução desta análise descritiva, foram selecionadas três ferramentas de IA generativa, cada uma representando uma abordagem arquitetônica e funcional distinta:

- ATLAS.ti: Trata-se de uma plataforma amplamente utilizada para Análise Qualitativa de Dados (QDA), que recentemente integrou funcionalidades de IA. Algumas dessas funcionalidades são potencializadas por modelos da OpenAI, que faz da ATLAS.ti uma ferramenta híbrida que combina recursos tradicionais de QDA com as capacidades dos LLMs modernos.

- ChatGPT (OpenAI): Representa uma das principais soluções de LLMs existente no mercado. Devido a sua versatilidade e popularidade, foi incluída entre as soluções estudadas para avaliar o desempenho do seu modelo de linguagem de propósito geral para a realização de uma tarefa de análise de pesquisa especializada.
- Manus: Consiste em uma plataforma de IA generativa que opera com base em um modelo de agentes autônomos. Diferentemente dos modelos tradicionais de prompt-resposta, como a solução do ChatGPT, essa arquitetura baseada em agentes permite a execução de tarefas complexas e multifacetadas de forma interativa e autônoma, representando uma das abordagens mais avançadas no campo da IA.

O material utilizado nesta análise constitui um subconjunto de uma coleção de dados mais ampla, gerada durante o segundo semestre letivo de 2023 para uma tese de doutorado. A pesquisa principal investiga a inter-relação entre a educação linguística em língua inglesa e o desenvolvimento da criticidade e agência discente. O dado específico submetido à análise pelas ferramentas de IA generativa foi a transcrição de uma entrevista semiestruturada, com duração aproximada de 123 minutos. A entrevista foi conduzida por um dos autores deste artigo com uma professora de língua inglesa do Instituto Federal do Espírito Santo. Para garantir a fidelidade ao processo de pesquisa original, a transcrição na íntegra foi utilizada, sem edições ou pré-processamento, replicando exatamente o material de análise da pesquisadora responsável pela tese. Por uma questão de privacidade, os nomes usados na transcrição são fictícios.

A seleção do corpus foi feita com base na priorização de dados analisados na tese em desenvolvimento. Como é comum em pesquisas de cunho etnográfico, sobretudo em um trabalho de doutorado, há um grande volume de dados gerados, uma vez que o pesquisador está imerso no campo investigado. Ao analisar os dados diretamente relacionados com o objetivo da tese, a pesquisadora ponderou que aqueles advindos da entrevista supracitada eram tangenciais ao foco do trabalho, o que possibilitou explorá-los neste estudo diante de outro enfoque investigativo. É importante enfatizar que a análise dos dados feita pelas diferentes ferramentas de IA generativa durante este processo de investigação não faz parte e não será utilizada em outro trabalho. Esses resultados foram obtidos única e exclusivamente para que o comportamento das diferentes IAs pudesse ser comparado e avaliado.

O objetivo central da proposta foi conduzir uma análise crítica das contribuições e apontamentos gerados por cada uma das três ferramentas de IA. Para tal, um prompt idêntico⁷ foi submetido a cada plataforma, instruindo a análise da transcrição da entrevista: identificar de que forma as perspectivas

⁷ Disponível em: <https://drive.google.com/file/d/19wCqEhYrtAqOGhJtV0EeHxnfoDpnXHS7/view?usp=sharing>.

apontadas pela professora, em suas respostas, estão em consonância com as teorias de Letramento Crítico. A formulação do prompt foi projetada para espelhar uma tarefa analítica autêntica no campo da Linguística Aplicada, com foco em um arcabouço teórico específico.

Além disso, buscamos entender como as IAs escolhidas fariam a análise fundamentadas em pressupostos que se distanciam do tecnicismo, do mecânico, do homogêneo ou padrão. O Letramento Crítico envolve muito mais do que o trabalho com determinado código linguístico, pois o foco incide na produção e na negociação de sentidos, na compreensão de língua como prática social, que é localizada em contextos sócio-históricos diversos.

Monte Mór (2014) explica que o LC teve influência da pedagogia freireana, uma vez que a questão social e a compreensão de privilégios por meio da educação está presente nessas propostas. A autora pondera que a sociedade contemporânea exige uma participação ativa e agentiva de cidadãos críticos, engajados socialmente, pois como aponta Ferrari (2018, p. 108), é preciso um letrar plural e diverso para que as pessoas possam “(...) discernir, desconfiar, problematizar as fontes, as representações, os discursos”.

Portanto, de posse das análises feitas pelas ferramentas, adotamos os seguintes procedimentos analíticos: leitura detalhada do material e elaboração de uma planilha para comparação dos trechos que foram analisados pelas IAs e das justificativas apresentadas. Assim, com base na análise de conteúdo, procedemos à comparação dos resultados fornecidos, com foco no conceito de Letramento Crítico empregado e nos autores citados pelas plataformas (embora não tenha sido solicitada a indicação de autores às ferramentas). Os documentos completos que guiaram a análise estão disponíveis por meio do link compartilhado no próximo item.

ANÁLISE DOS DADOS

As três IAs utilizadas nesse experimento ofereceram quantidades bem distintas de respostas para o comando ou prompt. Enquanto o ATLAS.ti e o ChatGPT apontaram um número aproximado de resultados, 7 e 13 respectivamente, a Manus apresentou 85 trechos da entrevista que estariam associados à teoria de letramentos críticos. A diferença na quantidade de respostas se repete quando consideramos a qualidade das justificativas contendo o embasamento teórico. Devido às semelhanças entre o ChatGPT e o ATLAS.ti, a análise e a discussão das respostas produzidas por essas duas ferramentas serão apresentadas conjuntamente. Posteriormente, o mesmo será feito com os resultados gerados pela Manus.

ATLAS.ti e ChatGPT

O ATLAS.ti e o ChatGPT apresentaram resultados bastante similares em números, extensão dos trechos, qualidade das justificativas teóricas e apresentação visual. Cinco excertos específicos aparecem nos resultados de ambas as ferramentas⁸. Os trechos, relativamente curtos, são recortes de falas mais longas. Para além desses cinco trechos, o ATLAS.ti apresentou outros dois, e o ChatGPT, oito⁹. Trechos ocultos, representados por links numéricos, também fazem parte das respostas do ATLAS.ti. Esses números direcionam o leitor para o arquivo com a entrevista, em que se pode visualizar o trecho completo ou uma outra fala que, teoricamente, tem um conteúdo parecido¹⁰. Como os trechos ocultos não trazem justificativas, eles não foram considerados nesta análise.

As justificativas apresentadas pelas IAs são bastante similares, sendo pautadas em palavras-chave comumente associadas ao Letramento Crítico, como consciência crítica, valorização da experiência dos estudantes, práticas sociais, desenvolvimento de cidadãos críticos, diversidade de perspectivas, dentre outras. Não há uma explicação elaborada acerca da teoria, mas uma abordagem superficial que requer um olhar atento e não simplista de apenas afirmar que “há” LC. Embora não tenha sido solicitado, o GPT cita alguns autores que estão, de certa forma, relacionados com a temática, mas que apresentam suas peculiaridades e, obviamente, essas nuances não são abordadas pelas IAs e cabe ao pesquisador conhecer quais são as obras essenciais em determinados temas.

Para a realização do experimento, não houve treinamento das IAs, isto é, as ferramentas não foram alimentadas pelos autores com textos nos quais as justificativas deveriam se basear. Isso tende a favorecer a ocorrência de vieses. Sem treinamento, as IAs difundidas atualmente reproduzem as premissas e escolhas realizadas pelos desenvolvedores e trabalhadores que estabelecem os bancos de dados utilizados para o treinamento. Como apontam Sampaio *et al.* (2024a), é notório que a produção norte-americana, europeia e anglo-saxã são as priorizadas pelos indexadores, o que demonstra domínio epistemológico de países hegemônicos.

Embora tenhamos gerado apenas uma pequena amostra, consideramos que as referências sugeridas pelo ChatGPT corroboram esse problema. Se considerarmos o total das respostas sugeridas pelo ChatGPT, percebemos que Freire foi citado quatro vezes, embora sua pedagogia crítica não seja um sinônimo do Letramento Crítico. Contudo, o Brasil tem vasta pesquisa sobre LC e um número

⁸ A tabela comparativa com os cinco trechos específicos está disponível em: <https://drive.google.com/file/d/1qo962RWNs0qIINW8r-e-q7FKNIkvHSCN/view?usp=sharing>.

⁹ Os resultados completos fornecidos pelo ChatGPT e pelo ATLAS.ti estão disponíveis em: https://drive.google.com/file/d/1jslh4DRNHN7NHzeu6W_4MJgKpiq_6e0T/view?usp=sharing.

¹⁰ O documento ATLAS.ti expandido, disponível em https://drive.google.com/file/d/1uXczdCcLmc35jPxrFqs-5vvNXLQTXh_Z/view?usp=share_link, mostra essa expansão.

considerável de pesquisadores e teóricos que trouxeram muitas contribuições para a área. Ainda assim, nenhum desses autores aparece nas referências do ChatGPT. Esse resultado comprova a ocorrência de vieses neste experimento, comportamento que era esperado.

Outro registro pertinente é o uso da expressão “mecânica da língua” em uma das justificativas do ATLAS.ti. Em “Este trecho reflete a visão de que o ensino deve ir além da mecânica da língua, promovendo um aprendizado que seja significativo e crítico, alinhando-se com os princípios do Letramento Crítico”, fica evidente que o significado pretendido é de um ensino que vá além das estruturas linguísticas, do ensino de gramática. Se considerarmos a definição de alucinação como um output que não parece correto (Lemos, 2024), pode-se afirmar que “mecânica da língua” é um exemplo dessa inconsistência apresentada pelo ATLAS.ti.

Outro exemplo de alucinação é a sugestão de um trecho que não está nos dados. “Trabalhamos temas como rede social, feminismo, questões raciais...”, que foi sugerido pelo ChatGPT, não se encontra na entrevista que foi oferecida à IA como material de análise. O trecho é, aparentemente, um resumo feito pela IA de uma fala mais extensa da professora:

(1) Excerto 1 – Recorte da entrevista original: 69:33 Eduarda: *e eu acho relevante a gente trabalhar nesse sentido. Eu penso que o ideal seria a gente realmente construir os nossos programas, as nossas propostas a partir do que vem deles, mas, por exemplo, na estruturação do curso, como a gente se organizou nesse segundo semestre, alguns temas foram vindo a partir do que foi sendo falado no primeiro tema. Vamos dizer assim, a gente começou falando sobre rede social, que é uma coisa que é a realidade deles. Esses meninos, a gente sabe disso, estão o tempo todo... eles vivenciam isso cada minuto da vida deles tudo está na rede social. Mas a partir dali você consegue observar se a turma tem uma pegada de discussão e que gostaria de falar de outros temas. Eu percebi logo que tinha uma galera ali da turma que era bastante engajada com relação ao feminismo, com relação às questões raciais então, foi encaixando. Na minha opinião, o ideal seria que a gente fosse construindo aos poucos e não ter coisas tão prontas. Embora pareça que está tudo conectado, eu acho que vai sendo construído. Eu fiz um pouco assim: então, na próxima aula dá pra levar tal coisa porque vai ficar legal vai encaixar.... vai fluir.*

Como pode ser observado no excerto 1, Eduarda está explicando para a pesquisadora a sua visão de ensino de línguas e como se deu a construção do curso e a escolha dos temas. Embora redes sociais, feminismo e questões raciais tenham sido mencionadas na fala, em momento algum a professora afirma ter trabalhado esses temas. Na verdade, com base em outros dados que foram gerados para a pesquisa mais ampla, sabe-se que feminismo não foi um dos temas usados em sala. Portanto, o ChatGPT modificou o texto original e gerou uma fala falsa da professora e uma informação incorreta, assunto que será retomado na próxima seção.

Com base nos resultados que obtivemos, consideramos que o ATLAS.ti e o ChatGPT apresentaram resultados insatisfatórios ao tentar identificar, por meio de análise de dados textuais, as formas em que as perspectivas de uma professora estão em consonância com uma visão de Letramento Crítico. As questões que justificam essa afirmação incluem as respostas com pouco embasamento teórico, a associação de alguns autores ao LC, os possíveis vieses e as alucinações.

Esse resultado adverso talvez possa ser explicado pelo formato do experimento: nossas análises foram feitas utilizando a primeira resposta obtida, isto é, avaliamos como as diferentes IAs reagiram ao prompt em um primeiro momento. Morgan (2023) sugere que o ChatGPT responde bem aos questionamentos feitos por humanos durante o processo de “análise” ou conversa com os dados. A partir de uma primeira resposta, a “matéria prima” vai sendo refinada a partir de outros prompts. Como nosso experimento não avançou nessa direção, devido ao escopo do estudo, nossos resultados podem, de certa forma, ser considerados limitados.

Tendo apresentado as ponderações acerca dessas ferramentas, passamos, a seguir, a discutir a análise de dados obtida com o uso da Manus.

Manus

A Manus traz uma análise bem mais ampla se comparada ao que foi apresentado pelo ChatGPT e pelo ATLAS.ti. Em linhas gerais, a Manus funciona de uma forma diferente. Como a versão gratuita da ferramenta tratou apenas uma fração dos dados (cerca de 14 minutos de entrevista), a conclusão da tarefa só foi possível com a aquisição de créditos. Assim, ao analisar o mesmo texto e executar a mesma tarefa, a Manus gerou 85 respostas das quais 26 foram descartadas por motivos variados: dois trechos se repetem e apresentam exatamente as mesmas análises; 17 são trechos repetidos e apontam basicamente para os mesmos aspectos do LC em suas justificativas; cinco trechos não são falas da professora Eduarda; os trechos 40 e 50 não fazem parte da entrevista. As respostas apresentadas pela Manus foram numeradas e estão disponíveis em material suplementar¹¹.

As análises fornecidas pela Manus não foram apresentadas de forma tabular. As respostas ao prompt são textos bem estruturados, bem fundamentados e, portanto, mais complexas quando comparadas às análises oferecidas pelo ChatGPT e pelo ATLAS.ti. A seleção de trechos, sua associação à teoria de Letramento Crítico, e a qualidade das justificativas apresentadas são coerentes com o processo de análise qualitativa de dados que seria feita por nós, pesquisadores. Várias premissas do Letramento Crítico estão presentes nas justificativas e muitas delas aparecem recorrentemente.

¹¹Disponível em: https://drive.google.com/file/d/1sLZEaklvWR1dTVfgj3dHjwxaG0TNazb/view?usp=share_link.

Embora algumas informações utilizadas nas respostas apresentem problemas que serão detalhados mais adiante no texto, a maioria das justificativas estão de acordo com o entendimento sobre Letramento Crítico abordado neste trabalho.

Considerando a qualidade textual e teórica dos resultados produzidos pela IA, há de se imaginar que essas respostas poderiam ser usadas para a produção de artigos, dissertações e teses. Isso nos remete ao debate que existe atualmente sobre a autoria dos produtos entregues por modelos de IA (Thornton, 2023), uma questão também discutida por Pack e Maloney (2023). Como se sabe, os grandes modelos de linguagem são treinados a partir de textos que foram criados por diversos autores. Assim, os conteúdos gerados por esses modelos derivam de materiais cujos autores originais não têm ciência de seu uso, muito menos concederam autorização para tal finalidade. Na falta de diretrizes e políticas que regulamentem o uso das IAs, bem como nossos preceitos sobre pesquisas em Linguística Aplicada, questionamos se o uso dessas respostas seria ético.

Para além da qualidade da maioria das respostas, um fator chama a atenção com relação aos resultados apresentados pela Manus: a ferramenta sugeriu tanto trechos referentes às práticas pedagógicas e visões de ensino da professora que envolvem o Letramento Crítico quanto trechos que mostram o Letramento Crítico na prática, isto é, como o LC se revela no dia a dia, nas vivências dos indivíduos, como pode ser observado no excerto 2:

(2) Excerto 2 – Recorte das respostas apresentadas pela Manus: 4 / *"...eu percebi também que não haveria nenhum desdobramento profissional daquilo que eu fiz academicamente, porque era um contexto muito limitado, né? Quem ficou bem ali naquela, ou quem conseguiu seguir, foram as pessoas que já tinham um respaldo de pai e mãe, que tinham dinheiro para montar uma agência, para colocar eles dentro de um veículo de comunicação influente, né? E eu não tinha nada disso, assim. Quem era eu?" (06:45 Eduarda)* / Este trecho é um exemplo claro de Letramento Crítico, pois a professora Eduarda analisa sua experiência acadêmica e profissional sob uma perspectiva socioeconômica. Ela reconhece que o sucesso em sua área anterior não dependia apenas de mérito individual, mas de privilégios e recursos financeiros (ter um respaldo de pai e mãe, que tinham dinheiro para montar uma agência, para colocar eles dentro de um veículo de comunicação influente). Isso demonstra uma consciência crítica das estruturas de poder e desigualdade que influenciam as oportunidades individuais. Essa análise se alinha com o Letramento Crítico, que busca capacitar os indivíduos a questionar e a desconstruir as relações de poder presentes nos textos e nos contextos sociais. A teoria de Paulo Freire sobre a leitura do mundo e da palavra, e a compreensão das estruturas sociais que as permeiam, é um forte embasamento para essa associação. /

A sugestão desse trecho e sua justificativa são bastante relevantes para esta análise. Ela mostra que, para responder ao comando, a IA verificou, no material analisado, não só as práticas pedagógicas

da professora, mas também suas atitudes e comportamentos na sua trajetória de vida. O excerto 2 é um exemplo de como a professora é letrada criticamente, uma vez que identifica e problematiza práticas sociais e relações de poder nas quais está inserida. Outros resultados que envolvem a percepção e o comportamento crítico da Eduarda incluem trechos em que ela versa sobre seu ingresso em uma universidade federal (item 5), sua vivência na universidade particular (item 8), a mercantilização da educação (item 20), sua experiência na escola pública (itens 22 e 23), a falta segurança pública (item 25), e a análise sobre a sua capacidade de dirigir um automóvel (item 26). A importância dessas respostas está no fato de que elas não haviam sido pensadas ou previstas pelos autores deste trabalho. São, portanto, elementos que emergiram dos dados devido à interação com a IA.

Resultados semelhantes foram observados por Hamilton *et al.* (2023) em um estudo que explorou o potencial do ChatGPT em gerar temas qualitativos a partir da transcrição de entrevistas. A análise consistia na geração de temas presentes nos dados, atividade que foi realizada tanto pelos pesquisadores quanto pelo ChatGPT. Os autores buscaram, então, estabelecer as diferenças e semelhanças entre essas análises e constataram que a IA identificou temas que não tinham sido notados por humanos, mas que, após reflexão dos pesquisadores, foram validados. Nossos dados corroboram os resultados de Hamilton *et al.* (2023) e, por isso, destacamos que a possibilidade da IA fazer apontamentos sobre dados não observados por humanos é um dos benefícios concretos do uso das ferramentas de IA na pesquisa qualitativa.

É importante também registrar a forma como Freire, única referência nas respostas da Manus, aparece nos resultados. Em quatro justificativas, a ferramenta associa as práticas ou comportamentos descritos à pedagogia Freireana (itens 4, 39, 42 e 52). Como mencionado anteriormente, embora o Letramento Crítico tenha a pedagogia crítica de Freire em suas bases, tratá-los como sinônimos pode reduzir a historicidade desses conceitos. No entanto, Duboc e Gattolin (2015) ponderam que, muitas vezes, a pedagogia crítica e o Letramento Crítico são usados de maneira intercambiável, pois a ideia de crítico está relacionada com as propostas Freireanas acerca da formação de cidadãos críticos e conscientes do contexto social e das relações de poder. É possível que, uma ferramenta que não tenha sido treinada faça associações mais diretas e menos subjetivas das teorias em questão.

Apesar da Manus ter surpreendido em relação à quantidade de trechos que foram associados ao LC e à qualidade das justificativas apresentadas, quatro tipos de problema foram identificados, sendo eles: descontextualização, inferências infundadas, citações falsas, e inversão da ordem das falas. As descontextualizações acontecem quando a justificativa é desenvolvida a partir de algo que, de fato, não faz parte dos dados. É como se a ferramenta de IA não tivesse um entendimento do que está sendo discutido entre a pesquisadora e a professora. O excerto 3 é um exemplo de descontextualização:

(3) Excerto 3 – Recorte das respostas apresentadas pela Manus: 53 / *"Gente não é para entender tudo." (95:20 Alice) e "Não era esse o objetivo, era mais que eles mesmos se soltassem." (96:12 Eduarda) / Este conjunto de falas, da entrevistadora e da professora, reforça um princípio fundamental do Letramento Crítico: a superação da compreensão literal e a valorização da agência do aluno. A ideia de que não é preciso entender cada palavra, mas sim o sentido geral e a capacidade de se expressar livremente, sem medo de correção, promove um ambiente de aprendizagem mais autêntico e menos punitivo. O Letramento Crítico busca desenvolver a capacidade de leitura do mundo, onde a fluência não se limita à decodificação, mas à construção de sentido e à participação ativa no processo comunicativo. A ênfase em "se soltassem" indica uma preocupação com a voz do aluno e sua autonomia na interação com a língua. /*

A resposta 53 é constituída por dois trechos separados por 52 segundos de entrevista. Nesse momento da geração de dados, discute-se o uso de textos autênticos nas aulas de línguas e, em seguida, a participação intensa de um grupo específico de alunos é trazida para a conversa. Esses alunos sempre pediam e eram autorizados a fazer a leitura em voz alta dos textos. Apesar dessa prática não fazer parte dos objetivos da professora, ela aceitava a participação para que eles se soltassem e se sentissem mais confortáveis em sala. Portanto, os dois trechos fazem referência a assuntos distintos que não são identificados pela ferramenta e, conseqüentemente, uma justificativa incorreta é gerada.

A professora não expressa, nesta fala, sua preocupação com a voz do aluno ou sua autonomia na interação com a língua, mas sim, uma valorização do interesse dos discentes pela oralidade e a participação deles nas aulas. Essa compreensão mais ampla da entrevista é viabilizada não só pelas informações que estão no texto, mas também pelo conhecimento que uma das autoras, pesquisadora que gerou os dados, tem desses participantes, do contexto, e das experiências vividas durante o semestre letivo. Isso reforça a importância do pesquisador como um instrumento de pesquisa (Paulus; Marone, 2025; Morgan, 2023; Bogdan; Biklen, 1994), portanto, não é possível terceirizar a análise de dados qualitativos. Estudo e conhecimento dos dados ainda é um requisito para que tenhamos confiabilidade na pesquisa qualitativa (Chubb, 2023).

A construção de várias justificativas foi considerada problemática devido às inferências que não se sustentam com base nos dados, ou trechos indicados. A resposta 6 (excerto 4) é um exemplo de como um resultado pode ser completamente equivocado: existe uma falta de compreensão sobre o que está sendo dito e as inferências vão muito além do que está disponível no trecho específico.

(4) Excerto 4 – Recorte das respostas apresentadas pela Manus: 6 / *"...a educação você pode até ganhar pouco. Mas você nunca vai deixar de ganhar. Se você, principalmente se você se identifica com essa, né? Com essa dinâmica de... De ensinar. De ensinar e tal." (11:50 Eduarda) / Este trecho, embora proferido pela sogra da professora Eduarda, é internalizado por ela e reflete uma visão de educação que*

transcende o aspecto financeiro, focando na identificação pessoal e na paixão pelo ensino. No Letramento Crítico, a educação é vista como um processo de transformação social e individual, e não apenas como um meio de ascensão econômica. A ideia de que "você nunca vai deixar de ganhar" se identifica com a noção de que o valor do ensino reside na sua capacidade de impactar vidas e promover o desenvolvimento humano, o que está em consonância com os princípios do Letramento Crítico que valorizam a formação de cidadãos conscientes e engajados. /

O aspecto financeiro era o ponto principal da fala. A professora estava realmente falando em ganhar dinheiro. Professores conseguem se inserir no mercado de trabalho facilmente e assim, mesmo que os salários sejam baixos, eles sempre encontram oportunidades de emprego. Além disso, o trecho não oferece elementos suficientes para a formulação da última oração, que foi desenvolvida a partir de uma compreensão incorreta do que foi posto pela professora. Na mesma linha, o excerto 5 apresenta um exemplo de como a ferramenta de IA desenvolve argumentos que não podem ser desencadeados a partir das ideias presentes na fala da professora.

(5) Excerto 5 – Recorte das respostas apresentadas pela Manus: 7 / *"Eu sou uma pessoa que eu gosto de trabalhar, eu gosto de dar aula. Então, assim, não teve nada a mais, entendeu? Então, você não precisa nem agradecer, assim."* (01:58 Eduarda) / *Embora pareça uma afirmação simples sobre gostar de trabalhar e dar aula, no contexto do Letramento Crítico, essa paixão pode ser interpretada como um motor para uma prática pedagógica engajada e transformadora. Professores que gostam do que fazem tendem a ir além do currículo prescrito, buscando formas mais significativas e relevantes de se conectar com os alunos e com o conteúdo. Essa dedicação intrínseca pode levar a uma abordagem mais crítica e reflexiva do ensino, onde o professor se sente motivado a questionar e a inovar, o que é fundamental para o desenvolvimento do Letramento Crítico nos alunos.*

A interpretação do item 7 é meramente especulativa. A paixão pelo ensino pode gerar vários desdobramentos, que não estão necessariamente vinculados a uma postura crítica de ensino. Apesar do texto apresentar uma sequência lógica, a argumentação construída atribuindo "uma abordagem mais crítica e reflexiva do ensino" à paixão pelo exercício da docência não se sustenta.

Um dos itens apresentados pelo ChatGPT foi considerado uma alucinação porque não era uma fala real da professora. Trata-se de um resumo de uma fala mais longa e contém informações incorretas. Diferentemente do ChatGPT, a Manus gerou dois trechos cujas origens não podem ser explicadas. A resposta 50¹² consiste em duas linhas que não têm sentido e estão supostamente localizadas no mesmo minuto da entrevista (104:42).

¹² A resposta 50 não foi considerada uma resposta válida por ser uma fala da pesquisadora. Entretanto, ela será usada nesta sessão como exemplo de alucinação da ferramenta.

(6) Excerto 6 – Resposta Manus seguido pelo trecho original da entrevista: 50 / *"...Eu não dirigi tanto quanto eu gostaria de ter dirigido, mas eu dirigi o suficiente para entender que o que a gente faz em sala de aula, às vezes, não é o que eles precisam lá fora."*(104:42 Pesquisadora)

104:42 Pesquisadora: *Mas eu acho que assim, ter um colega observando é um outro tipo de julgamento e aí eu acho que nem todo mundo...*

Eduarda: *Está aberto, né?*

A resposta 40 é um outro trecho que não faz parte da entrevista. Exceto pela primeira oração, o texto de cinco linhas sobre concepções de língua é bem construído e estaria supostamente localizado no minuto (82:51). Apesar de pertinente e alinhada com as concepções de Eduarda, a fala não é feita pela professora. No referido minuto, encontramos uma outra fala no texto original:

(7) Excerto 7 – Resposta Manus seguido pelo trecho original da entrevista: 40 / *"...Eu não dirigi tanto quanto eu queria. Mas eu acho que é sobretudo é a questão de que eu não queria que eles ficassem presos a uma regra. Eu queria que eles entendessem que a língua é viva, que a língua é fluida, que a língua é para se comunicar. E que eles podem usar a língua para se comunicar sobre o que eles quiserem. E que não tem certo e errado. Tem o que funciona e o que não funciona."*(82:51 Eduarda)

82:51 Eduarda: *Siiiiim, é muito isso, entendeu? Porque quando você começa a colocar regra, regra, regra, "e é assim", "só pode ser assim", "só tem esse gabarito e não vamos discutir outras respostas, tá gente?", "é isso"...*

83:07 Eduarda: *Você distancia, ou você tira a crença de que ele pode usar aquilo ali pra se expressar, entendeu? E o meu objetivo é exatamente o contrário, assim, tem sido. É o que eu tô te falando: nem sempre pode ser feliz esse objetivo, alcançado pra todo mundo. E, às vezes, os alunos têm a expectativa de que eu diga o que é certo e o que é errado. Eu sei disso, entendeu? A gente tá nesse lugar, né? A professora sabe se é certo ou se é errado.*

Registros sobre alucinações das IAs como a percepção de citações irreais feita por Pack e Maloney (2023) são comuns na literatura. Alkaissi e McFarlane (2023), por exemplo, relatam a ocorrência de alucinações quando os sistemas de IA são treinados em grandes quantidades de dados aleatórios. Contudo, os resultados gerados por este experimento revelam que as alucinações ocorrem mesmo em um conjunto muito pequeno de dados. O texto que foi carregado e submetido à análise pelas ferramentas de IA para a escrita deste artigo é consideravelmente curto quando comparado aos dados gerados em uma pesquisa qualitativa. Ainda assim, A Manus e o ChatGPT ofereceram trechos que não pertenciam à entrevista.

A identificação de duas dessas alucinações, assim como a compreensão do que está sendo discutido durante a entrevista, só foi possível devido ao conhecimento dos dados. Segundo Morgan

(2023), o conhecimento do pesquisador sobre os dados, que é fundamental para a construção de sentidos na pesquisa qualitativa, é também a chave para evitar respostas sem sentidos como as que foram observadas nos resultados gerados pelas ferramentas. A pesquisa qualitativa requer que o pesquisador estude cuidadosamente e conheça todo o conjunto de dados, mesmo quando as IAs sejam utilizadas na fase de análise textual (Morgan, 2023). O autor reitera algo que tem sido amplamente defendido: as IAs só podem funcionar como um auxílio, um fornecedor de insights no processo de análise de dados qualitativos (Paulus; Marone, 2025; Hamilton *et al.*, 2023).

A facilidade que as IAs prometem agregar à pesquisa qualitativa não chega sem prejuízos. Talvez o mais grave deles seja o distanciamento entre o pesquisador e os dados. Corroboramos a visão de Paulus e Marone (2025) ao pensar que a rapidez na geração de resultados em pesquisa qualitativa não está em consonância com suas bases epistemológicas. Acreditamos que a pesquisa requer tempo, leitura, releitura, anotações, estudo dos dados para uma melhor compreensão dos fenômenos, melhores reflexões e melhores resultados (Sampaio *et al.*, 2024b)

Por fim, foi observado que, nas respostas 44 e 78, a Manus inverteu a ordem em que as falas foram feitas.

(8) Excerto 8: Resposta da Manus e trecho original da entrevista: 44 / *"Eu não acho que isso é não ser profissional, porque depende da perspectiva. (...) Pois é, aí o que eu acho sobre aprendizado, ensinar nunca vai ser garantia de que se vai aprender. Então, a minha tentativa maior é que faça sentido para eles, quer seja o conteúdo ou a estrutura da língua, e que, de fato, desperte algum interesse com relação ao conteúdo e língua, muito mais do que que ele saiba que ele precisa usar a conjunção tal em tal lugar desse modo."* (80:46 Eduarda) /

80:09 Eduarda: *Pois é, aí o que eu acho sobre aprendizado, ensinar nunca vai ser garantia de que se vai aprender. Então, a minha tentativa maior é que faça sentido para eles, quer seja o conteúdo ou a estrutura da língua, e que, de fato, desperte algum interesse com relação ao conteúdo e língua, muito mais do que que ele saiba que ele precisa usar a conjunção tal em tal lugar desse modo.*

80:46 Eduarda: *Eu não acho que isso é não ser profissional, porque depende da perspectiva.*

No excerto 8, pode-se constatar que o trecho destacado acontece depois de "... em tal lugar desse modo". Essa é uma evidência de que a Manus alterou o texto original durante a execução do comando. A inserção de marcadores de tempo e a criação de falas, questões que já foram apontadas, corroboram essa constatação, que pode ser bastante problemática para a pesquisa qualitativa.

Como alertam Sinha *et al.* (2024), os mecanismos internos de funcionamento dos grandes modelos de linguagem ainda não são explícitos para o público em geral. Portanto, a incorporação dessas ferramentas na pesquisa qualitativa deve ser feita em meio a muitas considerações e cuidados. Lemos (2024) enfatiza que um erro produzido pela IA leva a um outro erro, como por exemplo, uma citação que

não existe pode gerar uma argumentação incorreta em um artigo. Nessa linha, as alucinações e as inversões de fala que foram observados neste experimento poderiam, em uma pesquisa real, gerar percepções de padrões, construções de sentido e conclusão irreais.

CONSIDERAÇÕES FINAIS

O objetivo deste trabalho foi explorar de que forma a IA generativa pode contribuir para a análise de dados qualitativos à luz de teorias da Linguística Aplicada. Primeiramente, reiteramos que, a nosso ver, desenvolver uma pesquisa qualitativa representa um processo que envolve arcabouços teóricos, metodologias de pesquisa, análise humana, interpretação de dados e, conseqüentemente, construção de sentidos (Paulus; Marone, 2025). Portanto, o experimento aqui descrito corresponde a uma das etapas do processo de análise de dados qualitativos, que é a análise textual dos dados.

Assim como o resultado de outras pesquisas acerca desse tema (Hamilton *et al.*, 2023; Prescott *et al.*, 2024), nossos dados indicam que as ferramentas de IA, no caso deste trabalho destacamos a Manus, oferecem a oportunidade de facilitar e melhorar a análise qualitativa de dados textuais se houver o acompanhamento e a revisão do pesquisador especialista no assunto, como também sugerido por Sampaio *et al.* (2024b) e Morgan (2023).

Entretanto, há diversos desafios que precisam ser considerados. Apesar da Manus ter indicado aspectos relevantes acerca da análise textual e ter apresentado coerência com o arcabouço teórico selecionado, há trechos que foram inventados pela ferramenta e certas inconsistências identificadas por meio da análise e interpretação humana, ou seja, pelos autores deste texto. As demais IAs utilizadas, ChatGPT e a ATLAS.ti, além de alucinações, apresentam análises superficiais que precisam ser aprofundadas.

Assim, ponderamos que a produção de conhecimento em Linguística Aplicada atualmente envolve o reconhecimento de questões éticas, culturais, sócio-históricas e ideológicas e suas implicações com as subjetividades do(a) pesquisador(a). Afinal, como alerta Tílio (2020, p. 31), “a pesquisa em Linguística Aplicada busca contribuir para o protagonismo de cidadãos críticos, capazes de promover transformação social por meio do conhecimento”. Portanto, é fundamental considerarmos se e como a IA generativa fará parte de procedimentos metodológicos envolvidos em pesquisas qualitativas, visto que seus usos implicam em posicionamentos ontoepistemológicos sobre a ética em pesquisa.

Esperamos que este experimento auxilie pesquisadores a identificar problemas que podem ser encontrados diante do uso de IA generativa, fomentando reflexões críticas e éticas sobre esse cenário. Estamos cientes da necessidade de discussões aprofundadas sobre esse tema e concluímos, com base

em Sampaio *et al.* (2024b), que o uso de IAs pode proporcionar atalhos, o que consequentemente pode implicar na produção de mais estudos, mas isso não significa que estamos, necessariamente, desenvolvendo pesquisas melhores. As IAs auxiliam na identificação e na justificativa da análise de dados, mas arriscamos dizer que ainda carecem de indagações e problematizações feitas por nós, até o presente momento, considerados humanos.

Referências

- ALKAISSI, H.; McFARLANE, S. I. Artificial hallucinations in ChatGPT: implications in scientific writing. *Cureus*, v. 15, n. 2, e35179, 2023. DOI: 10.7759/cureus.35179. Disponível em: <https://www.cureus.com/articles/138667-artificial-hallucinations-in-chatgpt-implications-in-scientific-writing>. Acesso em: 10 ago. 2025.
- BOGDAN, R.; BIKLEN, S. *Investigação qualitativa em educação*. Tradução: Maria João Alvarez; Sara Bahia dos Santos; Telmo Mourinho Batista. Porto: Porto Editora, 1994.
- CHUBB, L. A. Me and the machines: possibilities and pitfalls of using artificial intelligence for qualitative data analysis. *International Journal of Qualitative Methods*, v. 22, p.1-16, 2023. Disponível em: <https://journals.sagepub.com/doi/full/10.1177/16094069231193593>. Acesso em: 10 ago. 2025.
- CNPq. Portaria CNPq nº 2.664, de 6 de março de 2026. Institui a Política de Integridade na Atividade Científica do CNPq. Brasília: CNPq, 2026.
- D'ESPOSITO, M. E. W.; GATNER, S. AI in the Teaching-Learning of Languages. *The ESPECIALIST*, São Paulo, v. 45, n. 3, p. 134-153, 2024. DOI: 10.23925/2318-7115.2024v45i3e63941. Disponível em: <https://revistas.pucsp.br/index.php/esp/article/view/63941>. Acesso em: 30 jul. 2025.
- DUBOC, A. P.; GATTOLIN, S. *Letramentos e línguas estrangeiras: definições, desafios e possibilidades em curso*. In: FERRARETO, Rosana; LUCAS, Patrícia (Orgs.) *Temas e Rumos nas Pesquisas em Linguística (Aplicada): Questões empíricas, éticas e práticas*. Campinas: Pontes, 2015, p. 245-280.
- FERRARI, L. Letramento Crítico e formação de professores: uma conversa necessária. *PERCursos Linguísticos*, v. 8, n. 20, p. 105-116, 2018. Disponível em: <https://periodicos.ufes.br/percursos/article/view/21658>. Acesso em: 20 set. 2025.
- HAMILTON, L. *et al.* Exploring the use of AI in qualitative analysis: a comparative study of guaranteed income data. *International Journal of Qualitative Methods*, v. 22, p. 1-13, 2023. Disponível em: <https://journals.sagepub.com/doi/10.1177/16094069231201504>. Acesso em: 01 ago. 2025.
- LEMOS, A. L. Erros, falhas e perturbações digitais em alucinações das IA generativas: tipologia, premissas e epistemologia da comunicação. *MATRIZES*, v.18, n.1, p.75-91, 2024. Disponível em: https://revistas.usp.br/matrizes/pt_BR/article/view/210892. Acesso em: 12 ago. 2025.
- MARTINS, C.; PITZ, M. L. Escrevendo o futuro: perspectivas docentes sobre o uso de Inteligência Artificial na produção escrita em Língua Estrangeira. *Revista Horizontes de Linguística Aplicada*, v. 23, n.2, DT11, 2024. Disponível em: <https://doi.org/10.26512/rhla.v23i2.53707>. Acesso em: 30 jul. 2025.

- MONTE MÓR, W. Convergência e diversidade no ensino de línguas: expandindo visões sobre a "diferença". *Polifonia*, v. 21, n. 29, p. 234-253, 2014. Disponível em: <https://periodicoscientificos.ufmt.br/ojs/index.php/polifonia/article/view/1940>. Acesso em: 18 jul. 2025.
- MORGAN, D. Exploring the use of artificial intelligence for qualitative data analysis: the case of ChatGPT. *International Journal of Qualitative Methods*, v. 22, 1-10, 2023. Disponível em: <https://journals.sagepub.com/doi/10.1177/16094069231211248>. Acesso em: 10 ago. 2025.
- PACK, A.; MALONEY, J. Using generative artificial intelligence for language education research: insights from using OpenAI's ChatGPT. *TESOL QUARTERLY*, v. 57, n.4, p.1571-1582, 2023. Disponível em: <https://onlinelibrary.wiley.com/doi/10.1002/tesq.3253>. Acesso em: 10 ago. 2025.
- PAULUS, T.; MARONE, V. "In minutes instead of weeks": discursive constructions of Generative AI and qualitative data analysis. *Qualitative Inquiry*, v. 31, n. 5, p. 395-402, 2025. Disponível em: <https://journals.sagepub.com/doi/10.1177/10778004241250065>. Acesso em: 10 ago. 2025.
- PRESCOTT, M. *et al.* Comparing the Efficacy and Efficiency of Human and Generative AI: qualitative thematic analyses. *JMIR AI*, v. 3, 2024. Disponível em: <https://ai.jmir.org/2024/1/e54482>. Acesso em: 05 ago. 2025.
- SAMPAIO, R. *et al.* ChatGPT e outras IAs transformarão a pesquisa científica: reflexões sobre seus usos. *Revista de Sociologia e Política*, v.32, p.2-24, 2024a. Disponível em: <https://www.scielo.br/j/rsocp/a/rfSfWXpWqJWgrbRktcpXq9v/?format=pdf&lang=pt>. Acesso em: 10 ago. 2025.
- SAMPAIO, R. *et al.* Uma revisão de escopo assistida por inteligência artificial (IA) sobre usos emergentes de IA na pesquisa qualitativa e suas considerações éticas. *Revista Pesquisa Qualitativa*, v. 12, n. 30, p. 01-28, 2024b. Disponível em: <https://editora.sepq.org.br/rpq/article/view/729>. Acesso em: 25 jul. 2025.
- SILVA, P. R. B. S.; SANTOS JÚNIOR, G. P. dos (colab.). *Inteligência artificial, linguagens e educação*. Aracaju: Editora IFS, 2024. E-book. 117 p. ISBN 978-85-9591-221-2. Disponível em: <https://repositorio.ifs.edu.br/biblioteca/handle/123456789/2069>. Acesso em: 10 ago. 2025.
- SINHA, R. *et al.* *The role of generative AI in qualitative research: GPT-4's contributions to a grounded theory analysis*. Symposium on learning, design and technology. Delft, 2024. Disponível em: <https://dl.acm.org/doi/10.1145/3663433.3663456>. Acesso em: 10 ago. 2025.
- THORNTON, I. A special delivery by a fork: where does artificial intelligence come from? *New Directions for Evaluation*, p.23-32, 2023. Disponível em: <https://onlinelibrary.wiley.com/doi/abs/10.1002/ev.20560>. Acesso em: 10 ago. 2025.
- TÍLIO, R. C. 30 anos da ALAB: 30 anos de Linguística Aplicada e Ensino de Línguas no Brasil. *Raída*, v. 14, n. 36, p. 17 – 36, 2020. Disponível em: <https://ojs.ufgd.edu.br/index.php/Raído/article/view/12195>. Acesso em: 09 abr. 2026.
- VASWANI, A. *et al.* Attention is all you need. *Advances in neural information processing systems*, p. 5998-6008, 2017. Disponível em: <https://arxiv.org/abs/1706.03762>. Acesso em: 19 set. 2025.