

Tecnorrealismo e apropriação crítica: IA generativa no ensino de História na América Latina

Technorealism and Critical Appropriation: Generative AI in History Teaching in Latin America

Felipe Cittolin Abal

Doutor em História pela Universidade de Passo Fundo, Brasil
Professor do Programa de Pós-Graduação em História e do Programa de Pós-Graduação em Direito da Universidade de Passo Fundo, Brasil
felipe.c.abal@hotmail.com
<https://orcid.org/0000-0002-6208-5893>
<http://lattes.cnpq.br/2584975972936590>

Resumo: Este artigo investiga como educadores de História na América Latina podem apropriar-se criticamente de ferramentas de IA generativa. Partindo da análise das limitações técnicas dos modelos de linguagem e dos problemas do tecno-otimismo e do colonialismo digital, argumenta-se que recusar ou aceitar acriticamente essas tecnologias são posturas igualmente inadequadas. Propõe-se alternativa baseada no tecnorrealismo engajado: três modelos de GPTs personalizados (tutor socrático, aprendiz curioso e laboratório de fontes) desenhados para desenvolver pensamento histórico crítico. A proposta não resolve contradições estruturais do colonialismo digital, mas oferece caminhos de apropriação mediada que reduzem danos e ampliam capacidades críticas. Conclui-se pela necessidade de disputar sentidos e usos dessas tecnologias mediante mediação docente consciente, letramento digital crítico e transformações estruturais que incluam soberania tecnológica latino-americana.

Palavras-chaves: Inteligência Artificial; Ensino de História; Colonialismo Digital; América Latina; Educação.

Abstract: This article investigates how History educators in Latin America can critically appropriate generative AI tools. Based on analysis of technical limitations of language models and problems of techno-optimism and digital colonialism, it argues that refusing or uncritically accepting these technologies are equally inadequate positions. An alternative based on engaged technorealism is proposed: three models of customized GPTs (Socratic tutor, curious learner, and historical sources laboratory) designed to develop critical historical thinking. The proposal does not resolve structural contradictions of digital colonialism, but offers paths of mediated appropriation that reduce harm and expand critical capabilities. It concludes with the need to dispute meanings and uses of these technologies through conscious teacher mediation, critical digital literacy, and structural transformations including Latin American technological sovereignty.

Keywords: Artificial Intelligence; History Teaching; Digital Colonialism; Latin America; Education.

Introdução

O lançamento do ChatGPT em novembro de 2022 interrompeu o debate educacional global de forma abrupta. A ferramenta atingiu 100 milhões de usuários em dois meses, tornando-se o aplicativo de crescimento mais rápido da história. Desde então, sistemas como Claude, Gemini e Copilot passaram a integrar o cotidiano educacional, forçando instituições que ainda resistiam à adoção de ambientes virtuais a confrontar estudantes que já utilizam assistentes de IA para redigir trabalhos, responder questões e “dialogar” sobre conteúdos complexos. No Brasil e na América Latina, essa transição foi turbulenta, expondo fragilidades estruturais de sistemas educacionais já fragilizados.

A reação oscila entre extremos. Sam Altman, CEO da OpenAI, declarou que “a decolagem começou” e previu que “com a superinteligência, podemos fazer qualquer coisa” (ALTMAN, 2023). Esse entusiasmo corporativo permeia políticas públicas e produção acadêmica que trata IA como solução tecnológica para problemas estruturais da educação. Há, é claro, a posição oposta: críticos apontam plágio facilitado, aprendizagem superficial, substituição docente, amplificação de desigualdades e erosão da autonomia intelectual. Ambas as posições, no entanto, compartilham de uma premissa equivocada: tratam a tecnologia como variável independente, ignorando que o relevante não é o que a IA pode fazer, mas quem a utiliza, com quais objetivos, sob quais condições e com quais consequências para diferentes grupos sociais (SELWYN, 2021).

No ensino de História, essas questões assumem dimensão específica. História não é disciplina de transmissão de fatos estabelecidos, mas campo de construção interpretativa onde narrativas sobre o passado são elaboradas, contestadas e transformadas mediante análise rigorosa de evidências e debate argumentativo (PUREZA; ABAL, 2026). Formar estudantes no pensamento histórico exige desenvolvimento de capacidades de contextualização, análise crítica de fontes, compreensão de múltiplas perspectivas e construção de narrativas fundamentadas. Como sistemas de IA generativa, que operam por previsão estatística de sequências de palavras

sem compromisso com verdade, contexto ou rigor argumentativo, podem ser integrados a esse processo formativo?

A questão ganha urgência quando situada no contexto latino-americano. Como observam Couldry e Mejias (2019), vivemos em uma fase do colonialismo caracterizada pela extração contínua de dados produzidos no Sul Global, posteriormente processados e monetizados por plataformas sediadas no Norte. Quando instituições educacionais latino-americanas migram dados para plataformas das big techs, transferem não apenas informações administrativas, mas dados sobre processos de aprendizagem, produção intelectual e padrões comportamentais dos estudantes (AMADEU, 2020). A dependência de infraestrutura digital controlada por corporações estadunidenses, combinada com desigualdades estruturais de acesso tecnológico, transforma o uso de IA generativa no ensino de questão técnico-pedagógica em problema político de primeira ordem.

Essas assimetrias de poder não são externas ao funcionamento das IAs generativas. Sistemas como ChatGPT, Claude e Gemini são desenvolvidos por corporações estadunidenses, treinados predominantemente em inglês e com dados do Norte Global, ajustados mediante trabalho precarizado de revisores no Sul Global, gerando lucros apropriados privadamente enquanto externalizam custos ambientais massivos (O'DONNELL; CROWNHART, 2025). Seus vieses algorítmicos não são falhas técnicas a serem corrigidas, são expressões de assimetrias estruturais da ordem digital contemporânea.

Este artigo investiga como educadores de História na América Latina podem apropriar-se criticamente de ferramentas de IA generativa, reconhecendo potencialidades pedagógicas sem ignorar limitações técnicas e contradições políticas. Recusamos duas respostas simples: a aceitação acrítica que naturaliza problemas estruturais e a rejeição total que abdica de disputar sentidos e usos dessas tecnologias. Propomos, alternativamente, caminhos de apropriação mediada que reduzam danos, ampliem capacidades críticas de estudantes e docentes e mantenham o pensamento histórico como objetivo central. A pergunta que orienta este trabalho é: como educadores de História na América Latina podem apropriar-se criticamente de ferramentas de IA generativa, reconhecendo tanto suas potencialidades pedagógicas quanto suas limitações técnicas e contradições políticas?

O argumento se desenvolve em quatro partes. Primeiro, oferecemos um panorama histórico e técnico das IAs generativas, explicando como funcionam os *Large Language Models* e

porque é crucial compreender que esses sistemas apenas preveem sequências estatisticamente prováveis de palavras, sem “saber” verdadeiramente o que dizem. Segundo, desenvolvemos uma crítica ao tecno-otimismo, examinando como o solucionismo tecnológico (MOROZOV, 2013) desloca a atenção de problemas estruturais da educação e analisando vieses algorítmicos e custos humanos da produção de IA. Terceiro, aprofundamos essa análise através do conceito de colonialismo digital (COULDRY; MEJIAS, 2019), demonstrando como a dependência tecnológica latino-americana condiciona o uso educacional de IA. Quarto, apresentamos uma proposta prática: três modelos de GPTs personalizados para o ensino de História (tutor socrático, aprendiz curioso e laboratório de fontes históricas) inspirados no trabalho de E. Mollick e L. Mollick (2024), enfatizando que são ferramentas complementares, jamais substitutivas do trabalho docente. Por fim, explicitamos limites estruturais dessas propostas, identificando riscos persistentes e estabelecendo princípios éticos para uso responsável, reafirmando o papel insubstituível do professor (FORNASIER, 2021).

Este artigo contribui para literatura ainda escassa sobre uso de IA no ensino de Humanidades, particularmente no contexto latino-americano. Diferencia-se de abordagens puramente técnicas ao insistir que questões pedagógicas são inseparáveis de dimensões políticas, econômicas e epistemológicas. Também se diferencia de críticas puramente negativas ao apresentar propostas concretas de apropriação que, embora não resolvam contradições estruturais, buscam ampliar capacidades críticas dentro de condições existentes.

1. Das Máquinas de Turing aos Modelos de Linguagem: breve história da Inteligência Artificial

Quando estudantes pedem ao ChatGPT que escreva uma redação sobre a Revolução Francesa ou que explique o conceito de anacronismo histórico, poucos entendem o que realmente está acontecendo. Não há compreensão, pensamento ou conhecimento histórico sendo mobilizado. Há apenas um modelo estatístico prevendo quais palavras provavelmente virão a seguir em uma sequência textual. Compreender essa natureza fundamental dos *Large Language Models* exige retroceder às origens da inteligência artificial e entender como chegamos aos sistemas generativos que hoje permeiam a educação.

A história da inteligência artificial tem sua gênese em conceitos teóricos desenvolvidos no século XX, com papel central do matemático e cientista da computação Alan Turing. Já em 1936, Turing concebia uma máquina de computação digital abstrata, posteriormente denominada “máquina de Turing universal”, capaz de ser programada para realizar qualquer cálculo executável por um “computador humano”, uma pessoa operando sistematicamente com tempo e recursos ilimitados (COPELAND & PROUDFOOT, 2007). Isto estabeleceu as bases conceituais para o desenvolvimento posterior da computação e, consequentemente, da inteligência artificial.

A primeira menção explícita à inteligência computacional ocorreu em 1947, quando Turing proferiu palestra discutindo a possibilidade de máquinas agirem de forma inteligente, aprenderem com a experiência e superarem humanos no xadrez (COPELAND & PROUDFOOT, 2007). O trabalho de Turing durante a Segunda Guerra Mundial na quebra de códigos cifrados alemães também impulsionou o campo. Em 1948, ele elaborou um relatório não publicado intitulado “Maquinaria Inteligente”, considerado por muitos como o primeiro manifesto da IA, no qual refletia sobre a capacidade de máquinas aprenderem pela experiência (COPELAND & PROUDFOOT, 2007). O famoso “teste de Turing”, proposto em 1950 no artigo “*Computing Machinery and Intelligence*”, descrevia o “jogo da imitação”: um interrogador humano se comunicaria com dois jogadores, um humano e uma máquina, sem saber suas identidades, cabendo à máquina enganar o interrogador. Turing previa que no ano 2000 as máquinas passariam no teste 30% das vezes (MOLLICK, 2024). A previsão era ao mesmo tempo otimista e limitada. Otimista porque superestimava a velocidade do progresso tecnológico e limitada porque já naquele momento Turing reconhecia que a “inteligência” de máquinas seria diferente da humana.

Os primeiros programas de IA foram executados na Grã-Bretanha em 1951 e 1952, em parte devido à influência de Turing sobre a primeira geração de programadores. O termo “Inteligência Artificial” foi cunhado apenas em 1956, quando John McCarthy propôs um curso de verão na Universidade de Dartmouth com essa denominação, iniciando o chamado “primeiro verão das IAs” (WOOLDRIDGE, 2020). O entusiasmo era imenso e esse período foi marcado por avanços significativos. Allen Newell, Herbert Simon e J.C. Shaw criaram o “*Logic Theorist*”, o primeiro testador automático de teoremas, demonstrando que computadores podiam processar

símbolos, não apenas números (ERTEL, 2011). Essa descoberta era crucial porque aproximava máquinas do raciocínio abstrato humano, ou ao menos de uma simulação convincente dele.

Newell, Simon e Shaw prosseguiram desenvolvendo o *General Problem Solver* (GPS), cuja primeira versão operou em 1957. O GPS resolvia diversos quebra-cabeças através de tentativa e erro orientada por heurística, método que se tornaria central na IA. Outros programas iniciais notórios incluíam o Eliza e o Parry, que simulavam comportamento conversacional de um terapeuta humano e de uma pessoa paranoica, respectivamente. Ambos dependiam de respostas elaboradas antecipadamente e truques simples de programação, revelando as limitações da época (COPELAND & PROUDFOOT, 2007). O interessante é que mesmo esses programas rudimentares já conseguiam enganar usuários temporariamente, fenômeno que antecipa debates contemporâneos sobre a aparência de inteligência e inteligência genuína.

O período entre o início da década de 1970 e o fim dos anos 1980 ficou conhecido como “inverno das IAs”. A comunidade científica experimentava frustração com o pouco progresso em problemas centrais da área. As previsões otimistas não se concretizavam e as pesquisas no campo pareciam entrar em declínio terminal (WOOLDRIDGE, 2020). Essa crise não foi apenas técnica, mas conceitual. Os pesquisadores descobriam que replicar inteligência humana era mais complexo do que imaginavam. A resposta a essa crise foi o desenvolvimento dos sistemas especialistas, baseados na premissa de que os esforços anteriores focavam excessivamente em abordagens gerais. Esses sistemas eram programas projetados para resolver problemas em áreas limitadas, fundamentando-se no conhecimento de especialistas humanos previamente codificado. Embora tenham demonstrado utilidade prática em diagnósticos médicos e planejamento financeiro, seu auge terminou no final dos anos 1980, quando a realidade dos resultados não correspondeu ao entusiasmo inicial (WOOLDRIDGE, 2020). Novamente, o padrão se repetia. Promessas grandiosas seguidas de resultados modestos.

O foco da pesquisa deslocou-se então para a criação de agentes, sistemas de IA funcionando como entidades autônomas inseridas em ambientes específicos e encarregadas de realizar tarefas determinadas. A perspectiva de IA baseada em agentes, influenciada pela IA comportamental, ganhou força nos anos 1990 e tornou-se a ortodoxia da área (WOOLDRIDGE, 2020). Essa mudança representava reconhecimento importante, já que, ao invés de buscar inteligência geral, os pesquisadores concentravam-se em sistemas que executavam bem tarefas específicas em contextos delimitados.

Por volta de 2005, as atenções voltaram-se às redes neurais e ao aprendizado profundo. Redes neurais são modelos computacionais inspirados no funcionamento do cérebro humano, compostas por camadas de “neurônios artificiais” que processam informações, aprendem padrões e fazem previsões. O termo “aprendizado profundo” refere-se à profundidade dessas redes, com múltiplas camadas processando problemas em níveis diferentes. Camadas próximas à entrada lidam com conceitos de baixo nível, enquanto camadas mais profundas tratam de conceitos abstratos (WOOLDRIDGE, 2020). Em 2010, o lançamento da Siri para iPhone popularizou a interface baseada em agentes, com assistentes atuando cooperativamente com usuários. Essas tecnologias começavam a penetrar o cotidiano de milhões de pessoas, naturalizando a interação com sistemas de IA.

A aquisição da DeepMind pelo Google em 2014 trouxe a inteligência artificial novamente ao centro das atenções. O aprendizado de máquina tornou-se central, realizado através de redes neurais e aprendizado profundo. O objetivo do aprendizado de máquina é que o programa alcance resultados esperados sem receber a fórmula exata de como fazer isso. Essa capacidade fundamenta o funcionamento de modelos contemporâneos como o ChatGPT, que utiliza três tipos de aprendizado. No aprendizado supervisionado, o modelo é treinado com dados rotulados, nos quais cada entrada possui uma saída correspondente. Durante o treinamento, o modelo faz previsões e é corrigido quando erra. No aprendizado não supervisionado, o modelo trabalha com dados sem rótulos, tentando encontrar padrões autonomamente. No aprendizado por reforço, o agente aprende a tomar decisões interagindo com um ambiente e recebendo recompensas ou punições baseadas em suas ações.

O ChatGPT, por exemplo, é treinado com conjuntos massivos de textos de diversas fontes, expondo-o a grandes volumes para que aprenda padrões de linguagem e estruturas gramaticais. Sua tarefa fundamental é prever a próxima palavra em uma sequência textual com base nas anteriores. Após esse pré-treinamento, realiza-se um ajuste fino utilizando técnicas de aprendizado supervisionado ou por reforço para refinar suas capacidades. O ChatGPT e modelos similares pertencem à categoria de *Large Language Models* (LLMs), modelos de inteligência artificial treinados em quantidades massivas de texto para realizar tarefas relacionadas à linguagem natural, como entender, gerar, resumir e traduzir textos.

Contudo, é fundamental estabelecer desde agora uma compreensão realista das capacidades e limitações dessas tecnologias. No estado atual, ainda não estamos próximos de

uma inteligência artificial geral capaz de reproduzir plenamente as capacidades humanas. O entusiasmo gerado por empresas como a OpenAI, embora útil para valorizar ações no mercado financeiro, não corresponde aos resultados efetivos obtidos pelo uso dessas tecnologias. Os LLMs são modelos matemáticos generativos da distribuição estatística de tokens em vastos corpos de texto produzido por humanos. São generativos porque permitem interações, mas as questões que respondem são de tipo muito específico. “Aqui está um fragmento de texto. Diga-me como esse fragmento pode continuar. De acordo com seu modelo das estatísticas da linguagem humana, quais palavras provavelmente virão a seguir?” (SHANAHAN, 2024: 70). Essa é a operação fundamental, e tudo mais deriva dela.

Essa natureza estatística pode ser demonstrada através de exercício simples. Solicitar a um LLM que complete a frase “Penso, logo...” resultará na resposta “existo”, pois essa é a continuação estatisticamente mais provável. Em contraste, uma frase menos comum como “alienígenas invadiram São Paulo e trouxeram com eles...” produzirá respostas variadas, refletindo a ausência de padrão estatístico forte. Transformar um LLM em sistema de perguntas e respostas, integrando-o em sistemas maiores e utilizando engenharia de prompts para direcionar comportamentos, tornou-se o padrão atual. Mas isso não altera a natureza fundamental da operação.

Essa compreensão do funcionamento dos LLMs tem implicações para seu uso educacional, especialmente no ensino de História. Um LLM não possui crenças nem “sabe” verdadeiramente o que está respondendo. Embora as sequências possam parecer proposições, a conexão com a verdade é algo que apenas humanos percebem. O modelo não tem noção de verdade ou falsidade, pois não possui os mecanismos cognitivos para aplicar esses conceitos (SHANAHAN, 2024). Ajustar um modelo com feedback humano em larga escala, usando dados de preferências de avaliadores ou de uma base de usuários, é uma técnica que pode moldar respostas para refletir melhor normas de usuários, filtrar linguagem, melhorar precisão factual e reduzir tendência de gerar informações falsas. Contudo, isso não altera o cerne dos modelos, já que o resultado continua sendo um modelo da distribuição de tokens na linguagem humana, com modificações (SHANAHAN, 2024).

Atualmente, existe um senso comum fomentado por empresas de tecnologia e pessoas interessadas financeiramente de que a inteligência artificial é genuinamente inteligente e pode responder às mais variadas perguntas com precisão, apesar dos comuns erros e alucinações.

Esse senso comum é perigoso, particularmente na educação. Com estudantes utilizando cada vez mais essas ferramentas, torna-se necessário verificar de que modo os modelos de IA podem impactar o ensino, particularmente no campo da História, onde interpretação crítica, análise contextual e compreensão de múltiplas perspectivas constituem habilidades centrais. Se um sistema não compreende verdade ou falsidade, não contextualiza temporalmente, não distingue perspectivas históricas, como pode auxiliar na formação do pensamento histórico?

2. Entre promessas salvacionistas e realidades estruturais: uma crítica necessária

O debate atual sobre inteligência artificial na educação tem sido marcado por polarizações improdutivas. Encontramos, de um lado, o que Morozov (2013) denomina “solucionismo tecnológico”, a crença de que desafios educacionais complexos podem ser resolvidos mediante aplicação de ferramentas digitais, como se problemas estruturais fossem meras questões técnicas aguardando a solução algorítmica adequada.

Esse discurso messiânico, embora possa inspirar investimentos e inovação, opera uma perigosa naturalização da tecnologia como agente neutro e inevitável de progresso. É ignorado que a IA não opera no vácuo. Seus efeitos dependem radicalmente de quem a desenvolve, para quais finalidades, sob quais condições institucionais e com quais recursos materiais (SELWYN, 2021). No contexto educacional latino-americano, caracterizado por assimetrias de acesso e recursos, abraçar acriticamente essas promessas significa abdicar de enfrentar determinantes estruturais da desigualdade educacional. Como alerta Morozov (2013), o solucionismo tecnológico frequentemente substitui deliberação pública e enfrentamento político por promessas de automação, desviando atenção e recursos de transformações mais fundamentais.

Há, simultaneamente, resistências que tratam qualquer tecnologia digital como ameaça à autonomia docente e à qualidade do ensino. Essa postura, embora compreensível diante da precarização do trabalho educacional e da imposição vertical de “inovações” mal planejadas, também limita a capacidade de apropriação crítica dessas ferramentas. É necessário construir o que Selwyn (2021) denomina postura tecnorrealista, reconhecendo potencialidades sem ignorar

riscos, identificando condicionantes estruturais do uso tecnológico e insistindo que educadores e estudantes sejam sujeitos ativos no processo de integração tecnológica.

Essa postura exige compreender que os sistemas de IA carregam marcas profundas do mundo onde são construídos. Embora a IA possa parecer objetiva e neutra, os sistemas baseados em IA inevitavelmente incorporam vieses presentes em seus dados de treinamento, podendo reproduzi-los e amplificá-los sistematicamente (HAENLEIN; KAPLAN, 2019). Pesquisas demonstram, por exemplo, que sensores utilizados em veículos autônomos detectam com maior precisão tons de pele mais claros, resultado das imagens empregadas no treinamento dos algoritmos. Similarmente, sistemas de auxílio judicial podem reproduzir preconceitos raciais ao basear-se em histórico de julgamentos anteriores que refletem discriminações sociais preexistentes (HAENLEIN; KAPLAN, 2019).

No caso dos LLMs, a questão dos vieses torna-se mais complexa. Grande parte do treinamento desses modelos ocorre mediante retirada de dados da internet aberta, ambiente notoriamente permeado por discursos falsos, preconceituosos e desinformativos. Além disso, a seleção de quais dados serão coletados é definida primariamente por empresas, geralmente estadunidenses, cuja força de trabalho na computação é demograficamente homogênea e carrega seus vieses culturais, raciais e de gênero (MOLLICK, 2024). O resultado inevitável é uma visão distorcida do mundo, pois os dados utilizados estão longe de representar a diversidade da população que utiliza a internet e, por consequência, da população mundial.

À medida que a IA generativa torna-se mais utilizada na educação, esses vieses podem influenciar como estudantes percebem e interagem uns com os outros, bem como suas compreensões sobre estruturas sociais. No campo do ensino de História, esses vieses assumem dimensão adicional, podendo perpetuar narrativas hegemônicas sobre processos históricos, invisibilizar perspectivas marginalizadas e reforçar interpretações eurocêntricas de fenômenos globais.

Empresas de IA têm implementado estratégias para mitigar esses vieses. O DALL-E, por exemplo, adicionou discretamente a palavra “feminino” em número aleatório de solicitações para gerar imagens de “uma pessoa”, visando induzir diversidade de gênero não refletida nos dados de treinamento. A abordagem mais comum, contudo, é a intervenção humana através do Aprendizado por Reforço com Feedback Humano (RLHF), utilizado no ajuste fino dos LLMs (MOLLICK, 2024).

Esse processo de ajuste fino, contudo, carrega seus próprios problemas éticos. Trabalhadores mal remunerados de diversas partes do mundo, frequentemente do Sul Global, são contratados para ler e avaliar respostas geradas por IA, operando sob prazos curtos. Esses trabalhadores acabam expostos a todo tipo de conteúdo e alguns relatam terem sido traumatizados pela constante exposição a materiais gráficos e violentos que precisam revisar (MOLLICK, 2024). A ironia é cruel. Na tentativa de fazer com que as IAs atuem de forma ética, as empresas ultrapassam limites éticos com seus trabalhadores humanos.

Adicionalmente, o ajuste fino não é infalível. Os vieses em modelos avançados de LLMs tornaram-se mais sutis, mas não desapareceram (MOLLICK, 2024). No contexto educacional, essa persistência de vieses assume maior importância. Quais narrativas históricas são privilegiadas quando a IA é treinada predominantemente em inglês e com dados do Norte Global? Como processos históricos da América Latina são representados quando filtrados por algoritmos?

O uso de IA na educação apresenta ainda outros desafios que merecem atenção. O plágio automatizado torna-se facilitado. Ferramentas como LLMs permitem que textos gerados sejam entregues como se tivessem sido elaborados pelos estudantes, sem processo reflexivo ou autoria genuína, comprometendo tanto a formação quanto a integridade acadêmica. As alucinações representam outro problema sério, podendo induzir estudantes a erros graves. A UNESCO (2023) alerta que tais sistemas devem ser utilizados com supervisão rigorosa, respeitando princípios educacionais, para que não ampliem desigualdades e desinformação.

As discriminações algorítmicas constituem preocupação central. Silva (2025) argumenta que algoritmos carregam marcas do mundo onde são construídos, tendendo a reproduzir e até acentuar desigualdades raciais, de gênero e de classe. Os algoritmos operam por lógicas opacas, podendo reforçar processos de vigilância e exclusão. No campo educacional, isso pode significar exclusão de grupos com menor letramento digital, acesso limitado à tecnologia, ou cujas experiências não foram consideradas durante o treinamento.

Há também o custo ambiental, minimizado nos discursos tecno-otimistas. O treinamento e operação de modelos de IA consomem quantidades massivas de energia, gerando impactos globais através do uso intensivo de data centers alimentados predominantemente por fontes não renováveis. Ignorar esse impacto é incompatível com uma educação comprometida com justiça ambiental.

Diante desse panorama, torna-se evidente que o uso de IA na educação não pode ser tratado como mera questão de disponibilização de ferramentas. Exige letramento digital crítico de educadores e estudantes, compreensão das limitações técnicas e vieses estruturais desses sistemas e inserção em projetos pedagógicos que coloquem o pensamento crítico, a autonomia intelectual e a formação ética no centro. Como propõe a UNESCO (2023), o uso da inteligência artificial na educação deve ser centrado no ser humano, com supervisão adequada, respeitando diversidade e garantindo que a profissão docente seja valorizada. Isso implica não apenas disponibilizar ferramentas, mas compreender os limites e capacidades de estudantes em seus diversos aspectos.

3. Da extração de recursos à extração de dados: novas formas de subordinação

Os desafios identificados anteriormente constituem expressões de uma configuração estrutural mais profunda que caracteriza a ordem digital contemporânea: o colonialismo digital. Couldry e Mejias (2019) propõem esse conceito não como metáfora, mas como categoria que identifica continuidades históricas entre o colonialismo clássico e as formas contemporâneas de apropriação e exploração. Se o colonialismo histórico se concentrou na tomada de terras, recursos naturais e na exploração de corpos, o colonialismo digital opera através da extração contínua de dados produzidos por indivíduos e coletividades no Sul Global, transformando a experiência humana em matéria-prima a ser processada, monetizada e utilizada para fins de controle e acumulação.

Essa nova forma de colonialismo compartilha com seu predecessor histórico quatro características fundamentais: apropriação massiva de recursos (agora dados), criação de relações sociais profundamente assimétricas, produção de ideologias legitimadoras e concentração de riqueza e poder em atores do Norte Global (COULDRY; MEJIAS, 2019).

No contexto brasileiro e latino-americano, essas dinâmicas se manifestam de forma severa. A dependência tecnológica da região não resulta de déficit inevitável ou ausência de capacidades nacionais, mas de escolhas políticas que naturalizam e perpetuam a subordinação a corporações estrangeiras. Relatório produzido revela que entre 2014 e 2025, o Estado brasileiro

gastou aproximadamente R\$ 23 bilhões em contratos de softwares, serviços de nuvem e aplicações de segurança fornecidos por empresas estrangeiras. Desse montante, R\$ 10,3 bilhões foram gastos apenas nos anos de 2024 e 2025, acelerando a dependência. A Microsoft, isoladamente, recebeu mais de R\$ 3 bilhões em contratos públicos, enquanto Oracle, Google e Red Hat também figuram entre os principais beneficiários (SILVA *et al.*, 2025).

Esses números expressam mais que transferência de recursos financeiros. Consolidam uma arquitetura na qual o Estado brasileiro funciona como consumidor subordinado, adquirindo soluções controladas no exterior sem desenvolver capacidades próprias equivalentes. A expansão de datacenters no Brasil, frequentemente celebrada como modernização, não gerou autonomia tecnológica, pois a infraestrutura instalada permanece sob controle de corporações internacionais operando segundo lógicas de acumulação e governança definidas externamente (FURTADO; CUNHA, 2024).

O colonialismo digital articula-se economicamente através de padrões de dependência que replicam dinâmicas históricas de extração. A instalação de datacenters no país implica alto consumo energético e hídrico, frequentemente subsidiado pelo Estado, enquanto os lucros são repatriados para jurisdições onde se localizam as matrizes corporativas ou para paraísos fiscais. Essa transferência de valor é encoberta por retórica tecno-otimista que justifica investimentos mediante a invisibilização de impactos socioambientais e aprofundamento de assimetrias regionais (FURTADO; CUNHA, 2024).

A dimensão geopolítica dessa dependência caracteriza-se pela contínua extração de valor de dados produzidos no Sul Global, posteriormente processados e monetizados por plataformas sediadas no Norte. Estabelece-se uma relação profundamente assimétrica: países latino-americanos fornecem matéria-prima informacional e mercados consumidores, mas são excluídos das decisões sobre arquiteturas tecnológicas, padrões de governança e apropriação de excedentes econômicos (COULDRY; MEJIAS, 2019).

Juridicamente, a ausência de capacidade regulatória efetiva sobre plataformas transnacionais não decorre de impossibilidade técnica, são escolhas políticas que privilegiam autorregulação empresarial em detrimento de controle democrático. Mesmo legislações mais avançadas, como a Lei Geral de Proteção de Dados (LGPD), enfrentam limites de efetividade quando confrontadas com o poder econômico e lobby intensivo das big techs (POLIDO, 2024).

Essas dinâmicas ganham concretude particular no campo educacional. Quando sistemas de IA generativa são empregados em processos de ensino e aprendizagem na América Latina, reproduzem e, ao menos potencialmente, amplificam essas assimetrias estruturais.

O controle de dados educacionais torna-se problema estratégico. Quando instituições de ensino brasileiras e latino-americanas migram para plataformas das big techs, transferem-se não apenas dados administrativos, mas informações sobre processos de aprendizagem, produção intelectual de docentes e discentes, e padrões comportamentais de milhões de estudantes (AMADEU, 2020). Esses dados tornam-se matéria-prima para treinamento de algoritmos, refinamento de modelos de IA e desenvolvimento de produtos educacionais que serão posteriormente comercializados para essas mesmas instituições, completando ciclo de expropriação e mercantilização.

A barreira econômica reproduz desigualdades que estruturam nossos sistemas educacionais. Ferramentas mais avançadas de IA generativa são pagas e precificadas em dólares. Em contextos de desigualdade social profunda, isso significa que estudantes de classes privilegiadas acessam versões superiores dessas tecnologias, enquanto estudantes de classes populares dependem de versões gratuitas limitadas, quando possuem qualquer acesso. Não se trata de acaso, mas de reprodução de hierarquias sociais através de meios tecnológicos.

O letramento digital docente e discente enfrenta desafios que vão além do domínio técnico. Formar estudantes para uso crítico de IA exige não apenas compreensão de seu funcionamento, mas capacidade de analisar as dimensões políticas, econômicas e epistemológicas dessa tecnologia, uma análise que precisa ser situada nas condições específicas da América Latina. Como ensinar estudantes a questionar vieses algorítmicos se eles não compreendem as estruturas de poder que condicionam o desenvolvimento dessas tecnologias?

No ensino de História especificamente, essas questões ganham camadas adicionais. História não é mera acumulação de fatos, mas campo de disputas interpretativas onde narrativas sobre o passado são construídas, contestadas e transformadas. Quando IAs generativas são empregadas como ferramentas de apoio ao ensino histórico, carregam consigo visões de mundo, pressupostos e posicionamentos políticos incorporados em seus processos de treinamento e ajuste. Uma IA ajustada para evitar “controvérsias” pode omitir aspectos conflituosos de passados violentos. Uma IA que opera mediante lógica probabilística pode reforçar narrativas

hegemônicas, uma vez que essas são estatisticamente mais frequentes nos dados de treinamento.

Diante desse panorama, o uso de IA no ensino de História na América Latina não pode ser tratado como mera questão de disponibilização de ferramentas tecnológicas. Exige compreensão crítica das estruturas de poder que condicionam o desenvolvimento, controle e operação dessas tecnologias. Exige reconhecimento de que a soberania digital, entendida como capacidade de exercer controle sobre dados, infraestruturas e decisões tecnológicas que afetem direitos fundamentais e interesses nacionais, não será alcançada mediante soluções técnicas isoladas ou marcos regulatórios nacionais, mas através de estratégias que combinem desenvolvimento tecnológico, cooperação regional e contestação de estruturas hegemônicas vigentes (POLIDO, 2024).

É preciso, no entanto, distinguir dois níveis de análise que serão articulados na sequência. O nível estrutural atua na escala da política tecnológica do Estado, do investimento em pesquisa, da cooperação regional e da regulação democrática das plataformas. A soberania digital plena pertence a este nível e não pode ser obtida através de iniciativas pedagógicas individuais. O nível pedagógico, por sua vez, atua na escada da prática cotidiana, em condições dadas, com estudantes que já utilizam as ferramentas. As propostas apresentadas a seguir estão nesse segundo nível, não se tratando de caminho para a soberania digital, mas respostas que buscam formar usuários críticos dentro das estruturas vigentes. Tratá-las como solução para um problema estrutural seria reproduzir o solucionismo criticado. Por outro lado, toma-las como irrelevantes seria abdicar do trabalho formativo que cabe ao professor.

No contexto educacional, isso implica construir alternativas que não reproduzam acriticamente modelos importados e em desenvolver capacidades locais de criação e apropriação tecnológica, formar educadores capazes de mediação crítica e inserir uso de IA em projetos que coloquem autonomia intelectual, pensamento crítico e formação ética no centro. A próxima seção explora possibilidades concretas nessa direção, propondo modelos de uso de IA no ensino de História que, embora não resolvam as contradições estruturais identificadas, buscam reduzir danos e ampliar capacidades críticas.

4. Possibilidades de apropriação crítica: ferramentas customizadas e mediação docente

Diante do diagnóstico apresentado anteriormente seria possível concluir pela recusa total do uso de IA generativa na educação. Contudo, essa postura equivaleria a desistir de disputar os sentidos e usos dessas tecnologias, deixando-as inteiramente nas mãos de corporações e seus interesses mercantis. A alternativa que propomos não ignora as contradições, mas busca construir formas de apropriação que reduzam danos e ampliem capacidades críticas de estudantes e docentes. A ideia central é trabalhar com GPTs personalizados: ferramentas customizadas pelo próprio professor, alimentadas com materiais de qualidade selecionados e inseridas em projetos pedagógicos que mantêm o pensamento histórico crítico como objetivo fundamental.

A personalização de modelos de linguagem oferece vantagens sobre o uso direto de IAs genéricas. Permite ao docente controlar quais fontes alimentam as respostas do sistema, privilegiando historiografia academicamente reconhecida em detrimento a conteúdos da internet aberta. Além disso, possibilita configurar comportamentos pedagógicos específicos. Ao invés de simplesmente fornecer respostas prontas, o GPT pode ser instruído a fazer perguntas, estimular contextualização, apontar múltiplas interpretações. O modelo opera primariamente sobre o conteúdo limitado e verificado fornecido pelo professor, reduzindo (mas não eliminando) o problema das alucinações. Mais importante, o professor é mantido como protagonista do processo. É ele quem define objetivos de aprendizagem, seleciona materiais, desenha interações pedagógicas e avalia resultados (MOLLICK, E.; MOLLICK, L., 2024).

Inspirados no trabalho de E. Mollick e L. Mollick (2024) sobre inovação pedagógica com IA, propomos três modelos de GPTs personalizados para o ensino de História, cada um com função pedagógica distinta e adequado a momentos diferentes do processo de aprendizagem. Trata-se de propostas formuladas a partir de literatura disponível, ainda não submetidas a teste em sala de aula. Sua apresentação cumpre a função de oferecer modelos de experimentação, crítica e refinamento por outros educadores, bem como base para investigação empírica posterior. O primeiro modelo opera como tutor socrático, o segundo inverte a lógica tradicional

fazendo o estudante ensinar e o terceiro cria simulações que desenvolvem habilidades de análise de fontes primárias.

Nenhum desses modelos substitui a leitura de textos historiográficos, a análise de documentos originais, ou a discussão coletiva em sala de aula. Funcionam como ferramentas complementares que podem ampliar oportunidades de aprendizagem, especialmente em contextos em que estudantes dispõem de pouco tempo para atendimento individualizado.

O primeiro modelo, que denominamos “Tutor Socrático de História”, destina-se a auxiliar estudantes em momentos de estudo autônomo, oferecendo suporte quando o professor não está disponível. Diferentemente de um GPT genérico que só responde perguntas, o tutor foi instruído a nunca fornecer respostas prontas. Quando um estudante pergunta, por exemplo, “O que causou a Revolução Francesa?”, ao invés de listar fatores, o tutor responde, por exemplo: “Interessante pergunta! Antes de explorarmos as causas, você já estudou algo sobre a França no século XVIII? Como era organizada a sociedade francesa naquele período?”. O objetivo é estimular o estudante a construir conhecimento ativamente, mobilizando o que sabe, fazendo conexões e formulando hipóteses.

O tutor deve ser alimentado com obras historiográficas selecionadas sobre os temas do curso (digamos, textos de Eric Hobsbawm sobre revoluções) e instruído a basear suas orientações exclusivamente nesses materiais. Quando detecta equívocos ou simplificações, não corrige diretamente, mas faz novas perguntas que conduzem o estudante a perceber problemas em seu raciocínio. Por exemplo, se um estudante afirma que “a abolição foi um ato de bondade da Princesa Isabel” o tutor pode responder: “Entendo. Mas me diga: havia movimentos sociais lutando pela abolição naquele período? Qual era a situação econômica da escravidão em 1888? Como outros países haviam lidado com essa questão?”

Essa abordagem desenvolve autonomia intelectual e pensamento crítico, ao invés de perpetuar relação passiva com o conhecimento. Um exemplo de prompt técnico para configurar esse GPT seria:

Você é um tutor de História para ensino médio. Seu objetivo é desenvolver pensamento histórico crítico nos estudantes através de perguntas orientadoras. Nunca forneça respostas prontas ou em formato de lista. Quando um estudante faz uma pergunta, primeiro investigue o que ele já sabe sobre o tema. Faça

perguntas que estimulem contextualização temporal e espacial. Aponte sempre para múltiplas interpretações historiográficas quando relevante. Utilize apenas os documentos fornecidos como base de conhecimento. Ao detectar erros factuais ou interpretações simplificadas, não corrija diretamente, faça perguntas que levem o estudante a repensar sua posição.

O segundo modelo inverte a lógica pedagógica tradicional: ao invés do estudante aprender com o GPT, é o estudante quem ensina ao GPT. Pesquisas em ciências da aprendizagem demonstram que ensinar é uma das formas mais eficazes de aprender, pois exige organização do conhecimento, identificação de conceitos centrais, antecipação de dúvidas, e formulação clara de explicações (MOLLICK, E.; MOLLICK, L., 2024). Este modelo, que chamamos “O Aprendiz Curioso”, simula um estudante que está começando a estudar determinado tema e precisa da ajuda do usuário para compreendê-lo.

Podemos citar, por exemplo, um GPT configurado para aprender sobre o processo de independência do Brasil. Quando o usuário inicia a interação, o GPT se apresenta: “Oi! Estou estudando a independência do Brasil. Você pode me ajudar a entender? Ouvi dizer que tem a ver com D. Pedro e o grito do Ipiranga, mas não sei direito o que aconteceu”. O estudante, então, precisa explicar, em suas próprias palavras, o processo histórico. O GPT reage com entusiasmo quando as explicações são claras e precisas. Quando há imprecisões ou lacunas, expressa dúvidas de forma gentil: “Hmm, você disse que foi pacífico, mas será que não houve conflitos? Ouvi falar de uma guerra na Bahia...”.

Esse modelo é especialmente útil para fixação de conteúdo após aulas expositivas. O estudante já teve contato com o tema, leu textos ou assistiu aulas e agora precisa consolidar esse conhecimento explicando. O GPT foi alimentado com um texto específico sobre o tema (digamos, capítulo de livro didático sobre independência) e conhece as informações que o estudante deveria dominar. Ao final da interação, oferece feedback breve: “Você explicou muito bem a questão política e o papel das elites, mas poderia desenvolver mais a dimensão social – como isso afetou escravizados, indígenas, pessoas pobres livres?”. Esse feedback não substitui a avaliação docente, mas oferece retorno imediato que pode orientar estudos posteriores.

O terceiro modelo cria simulações gamificadas que desenvolvem habilidades do pensamento histórico: análise de fontes primárias, contextualização, identificação de

perspectivas e vieses, construção de narrativas baseadas em evidências. Denominamos esse modelo “Laboratório de Fontes Históricas”. Diferentemente dos dois anteriores, este não simula diálogo pedagógico convencional, mas cria situações-problema onde o estudante precisa agir como historiador, trabalhando com documentos de época.

Uma possível implementação: o GPT apresenta ao estudante uma carta escrita por um colonizador português no século XVI descrevendo populações indígenas. Após o estudante ler a carta, o GPT faz perguntas: “Quem escreveu este documento? Para quem? Com que finalidade? Que imagem dos indígenas o autor constrói? Podemos confiar plenamente nesta fonte?”. O estudante precisa responder, demonstrando capacidade de ler criticamente a fonte, identificar seu lugar de produção, reconhecer seus limites e vieses e contextualizá-la historicamente.

A implementação prática desses GPTs exige planejamento e mediação docente constante. A preparação começa com seleção de fontes de qualidade e criação do GPT personalizado, seguida de testes para identificar limitações que devem ser compartilhadas com estudantes. O letramento digital crítico precisa acontecer presencialmente antes do uso. Explicar que LLMs são modelos estatísticos que preveem sequências de palavras sem “saber” realmente o que dizem, demonstrar casos de alucinações e vieses, ensinar técnicas de prompts e discutir questões éticas como plágio e citação adequada.

O uso mediado propriamente dito disponibiliza os GPTs para atividades estruturadas. Por exemplo: ler capítulos sobre Revolução Industrial, depois interagir com o Tutor Socrático sobre dúvidas específicas, registrando perguntas e respostas para discussão coletiva posterior, ou ensinar ao Aprendiz Curioso sobre capitanias hereditárias e produzir parágrafo reflexivo sobre o próprio processo de aprendizagem. Sempre deve haver produção reflexiva associada. Oferecer alternativas para estudantes que não possam ou não queiram usar IA também é fundamental.

A avaliação crítica coletiva encerra o ciclo: discussão presencial sobre onde o GPT ajudou efetivamente, onde errou ou simplificou, como suas respostas se comparam aos textos lidos e que habilidades foram ou não desenvolvidas. Essa reflexão é essencial para formar estudantes que compreendem possibilidades e limites dessas tecnologias.

No contexto latino-americano, a implementação dessa proposta enfrenta desafios que não podem ser ignorados. A questão do acesso é fundamental. Ferramentas como ChatGPT têm versão gratuita limitada e versão paga necessária para criar GPTs personalizados, colocando dilemas sobre quem paga (professor, escola, estudantes) e como lidar com limites de mensagens.

A infraestrutura é outro obstáculo: internet instável e ausência de dispositivos adequados exigem soluções como laboratórios escolares ou uso coletivo, todas com limites práticos.

O letramento digital de docentes e discentes frequentemente é insuficiente para uso crítico dessas tecnologias, exigindo investimento em formação continuada, tempo institucional para experimentação e recursos para materiais. Além disso, persistem as questões estruturais. Mesmo com customização, os GPTs carregam vieses, consomem energia massiva e geram lucros para corporações estrangeiras. Usar essas ferramentas criticamente não resolve o colonialismo digital, apenas mitiga alguns de seus efeitos mais imediatos.

Diante dessas limitações, precisamos explicitar o que essa proposta não faz e não pretende fazer. Não resolve desigualdades estruturais de acesso à tecnologia, não substitui a necessidade de desenvolvimento de alternativas tecnológicas latino-americanas, não elimina os problemas éticos e ambientais da IA e não transforma automaticamente estudantes em pensadores críticos. O que a proposta oferece são caminhos possíveis para apropriação crítica de tecnologias que, querendo ou não, já estão presentes no cotidiano educacional. Melhor disputar seus usos e sentidos através de mediação docente consciente do que abandoná-las à lógica mercantil ou proibi-las de forma que serão usadas às escondidas, sem qualquer orientação.

Considerações Finais

Este artigo partiu da seguinte questão: como educadores de História na América Latina podem apropriar-se criticamente de ferramentas de IA generativa? A resposta é complexa. Não há solução técnica para problemas que são políticos, pedagógicos e epistemológicos ou apropriação possível que ignore as estruturas de poder que condicionam o desenvolvimento, controle e operação dessas tecnologias. Mas também não há caminho na recusa paralisante ou na proibição que simplesmente empurra o uso para a clandestinidade, sem orientação crítica.

Demonstramos que IAs generativas são sistemas de previsão estatística de sequências de palavras, não entidades que “sabem” ou “compreendem” o que dizem. Antropomorfizar esses sistemas, atribuindo a eles inteligência, consciência ou capacidade de raciocínio, leva a usos acríticos e confiança injustificada em suas respostas.

Argumentamos também que o tecno-otimismo que permeia discursos corporativos e produções acadêmicas sobre IA na educação constitui o que Morozov (2013) denominou de solucionismo tecnológico: a crença de que problemas sociais complexos podem ser resolvidos mediante aplicação de tecnologias digitais, deslocando atenção de causas estruturais. Desigualdades educacionais na América Latina não decorrem da falta de ferramentas tecnológicas sofisticadas, mas de precarização docente, infraestrutura inadequada, desigualdade social profunda e subordinação histórica. Introduzir IAs generativas nesses contextos sem transformar condições estruturais tende a amplificar desigualdades preexistentes.

Demonstramos, ainda, que o uso educacional de IA não pode ser dissociado das dinâmicas de colonialismo digital que caracterizam a ordem contemporânea. Quando instituições brasileiras e latino-americanas utilizam LLMs, não estão adotando ferramentas neutras, estão se inserindo em relações assimétricas de extração de dados, transferência de valor e subordinação tecnológica. Cada interação alimenta sistemas treinados majoritariamente no Norte Global, controlados por corporações estrangeiras, geradores de lucros apropriados privadamente enquanto custos ambientais e sociais são externalizados para o Sul.

Diante desse diagnóstico crítico, apresentamos proposta de apropriação mediada: três modelos de GPTs personalizados desenhados para desenvolver pensamento crítico ao invés de simplesmente fornecer respostas prontas. Inspirados no trabalho de E. Mollick e L. Mollick (2024), esses modelos são customizados pelo professor, alimentados com obras de qualidade e integrados a projetos pedagógicos que mantêm leitura de textos, análise de fontes primárias e debate coletivo como elementos centrais. Todo o processo envolve a observância dos princípios expostos.

Insistimos nos limites estruturais dessa proposta. GPTs personalizados não eliminam vieses algorítmicos, não impedem alucinações, não resolvem problemas de plágio, não substituem dimensões fundamentais do pensamento histórico e não transformam as condições de colonialismo digital que condicionam seu funcionamento. O que oferecem são caminhos possíveis de redução de danos e ampliação de capacidades críticas dentro de condições existentes.

A postura que defendemos, portanto, não é de otimismo ingênuo nem de pessimismo paralisante, mas de tecnorrealismo engajado. Cabe a educadores, estudantes e comunidades

educacionais disputar sentidos e usos dessas tecnologias, ao invés de abandoná-las inteiramente à lógica mercantil ou aceitá-las acriticamente em seus termos.

Essa disputa, contudo, não pode ser travada apenas no nível micro das práticas pedagógicas individuais. Exige também transformações estruturais que vão muito além do que este artigo propõe. O desenvolvimento de capacidades tecnológicas latino-americanas demanda investimento massivo. A soberania digital não será alcançada mediante compras de soluções do exterior, mas por meio formação de profissionais, financiamento de pesquisa e desenvolvimento, criação de infraestruturas públicas e cooperação regional. Universidades públicas latino-americanas, em particular, deveriam estar na vanguarda desse processo, não subordinadas a plataformas corporativas estrangeiras.

Marcos regulatórios efetivos que subordinem tecnologias digitais a controle democrático são igualmente necessários. Políticas públicas de formação docente precisam ir além do treinamento instrumental no uso de ferramentas, investindo em compreensão crítica das dimensões políticas e epistemológicas das tecnologias digitais. Simultaneamente, condições dignas de trabalho docente são incontornáveis. Professores com cargas horárias excessivas, salários inadequados, turmas superlotadas e infraestrutura precária não têm condições materiais de implementar propostas, por mais bem-intencionadas que sejam. A precarização do trabalho docente é barreira a qualquer inovação genuína, com ou sem tecnologia. Valorizar profissionalmente educadores não é custo, é investimento em projeto de desenvolvimento soberano.

A democratização real de acesso tecnológico completa esse quadro de transformações necessárias. Enquanto estudantes de classes privilegiadas acessam versões premium de IAs generativas, laboratórios bem equipados e internet de alta velocidade, estudantes de classes populares dependem de versões gratuitas limitadas, conexões instáveis e dispositivos inadequados. Políticas de inclusão digital precisam ir além de distribuir tablets ou laptops, fornecendo infraestrutura de conectividade universal, formação adequada e garantia de acesso a ferramentas de qualidade.

É central reafirmar o que consideramos fundamental: nenhuma tecnologia substitui o papel formativo do professor. Educar não é transmitir informações, mas formar sujeitos autônomos, críticos, éticos, capazes de pensamento rigoroso, sensíveis a complexidades, dispostos a questionar certezas, comprometidos com justiça. Isso exige relação humana,

presença, escuta, cuidado, dimensões que não podem ser reduzidas a algoritmos. Como observa Fornasier (2021), transformação educacional necessária não é técnica, é epistemológica: colocar formação crítica, ética e prospectiva no centro, em contraposição a modelo meramente transmissivo.

Os três modelos propostos aguardam verificação empírica em contextos de ensino e seus resultados efetivos só poderão ser avaliados após a implementação e estudo de seus impactos sobre a aprendizagem. Este artigo buscou contribuir para essas lutas oferecendo análise crítica das contradições em jogo e propostas concretas de apropriação mediada. Esperamos que estimule debates, inspire experimentações e seja questionado, testado e adaptado por educadores em diferentes contextos. O futuro da educação histórica não deve ser determinado por algoritmos desenvolvidos no Vale do Silício, mas por escolhas coletivas de educadores, estudantes e sociedades comprometidas com projeto educativo democrático, crítico e soberano.

Referências bibliográficas

- ALTMAN, Sam (2023). Reflections. *Blog pessoal*. Disponível em: <<https://blog.samaltman.com/reflections>>. Acesso em: 24 jun. 2025.
- AMADEU, Sérgio (2020). *Brasil, colônia digital*. Instituto Humanitas Unisinos. Disponível em: <<https://ihu.unisinos.br/categorias/600360-brasil-colonia-digital-artigo-de-sergio-amadeu>>. Acesso em: 6 jan. 2026.
- COPELAND, Jack B. & PROUDFOOT, Diane (2007). Artificial Intelligence: History, Foundations, and Philosophical Issues. In: THAGARD, Paul. *Handbook of the Philosophy of Psychology and Cognitive Science*. Amsterdam: Elsevier.
- COULDRY, Nick & MEJIAS, Ulises A. (2019). *The Costs of Connection: How Data Is Colonizing Human Life and Appropriating It for Capitalism*. Stanford: Stanford University Press.
- ERTEL, Wolfgang (2011). *Introduction to Artificial Intelligence*. London: Springer London. (Undergraduate Topics in Computer Science).
- FORNASIER, Mateus de Oliveira (2021). Ensino Jurídico no Século XXI e a Inteligência Artificial. *Revista Opinião Jurídica*, Fortaleza, ano 19, n. 31, pp. 1-32, maio/ago.
- FURTADO, Renato Guimarães & CUNHA, Simone Evangelista (2024). Inteligência Artificial, Data Centers e Colonialismo Digital: impactos socioambientais e geopolíticos a partir do Sul Global. *Liinc em Revista*, Rio de Janeiro, vol. 20, n. 2.
- HAENLEIN, Michael & KAPLAN, Andreas (2019). A Brief History of Artificial Intelligence: On the Past, Present, and Future of Artificial Intelligence. *California Management Review*. California, vol. 61, n. 4, pp. 5-14.
- MOLLICK, Ethan (2024). *Co-intelligence: living and working with AI*. New York: Portfolio/Penguin.

- MOLLICK, Ethan & MOLLICK, Lilach (2024). Instructors as Innovators: A future-focused approach to new AI learning opportunities, with prompts. *The Wharton School Research Paper*. Disponível em: <<https://ssrn.com/abstract=4802463>>. Acesso em: 4 jul. 2025.
- MOROZOV, Evgeny (2013). *To Save Everything, Click Here: The Folly of Technological Solutionism*. New York: PublicAffairs.
- O'DONNELL, James & CROWNHART, Casey (2025). We did the math on AI's energy footprint. Here's the story you haven't heard. *MIT Technology Review*. 20 mai. Disponível em: <<https://www.technologyreview.com/2025/05/20/1116327/ai-energy-usage-climate-footprint-big-tech/>>. Acesso em: 2 jul. 2025.
- POLIDO, Fabrício Bertini Pasquot (2024). Estado, soberania digital e tecnologias emergentes: interações entre direito internacional, segurança cibernética e inteligência artificial. *Revista de Ciências do Estado*. Belo Horizonte, vol. 9, n. 1.
- PUREZA, Fernando Cauduro; ABAL, Felipe Cittolin. Inteligência artificial no ensino de História: ChatGPT e ditadura militar brasileira. *Estudos Ibero-Americanos*, Porto Alegre, v. 52, n. 1, pp. 1-16, jan./dez. 2026.
- SELWYN, Neil (2021). *Education and Technology: Key Issues and Debates*. 2. ed. London: Bloomsbury Academic.
- SHANAHAN, Murray (2024). Talking about Large Language Models. *Communications of the ACM*. New York, vol. 67, n. 2, pp. 68–79.
- SILVA, Ergon Cugler de Moraes; ROCHA, Isabela; VAZ, José Carlos; VENEZIANI, Julia Ribeiro de Almeida & MODANEZ, Camila de Camargo (2025). *Contratos, Códigos e Controle: A Influência das Big Techs no Estado Brasileiro*. São Paulo: GETIP/EACH/USP; Brasília: GEPSI/IREL-UnB, jul. Disponível em: <<https://bit.ly/contratos-big-techs>>. Acesso em: 27 dez. 2025.
- SILVA, Tarcízio (2025). *Racismo Algorítmico e Regulação de Inteligência Artificial: O Contrato Racial na Produção do PL 2338/2023*. Tese (Doutorado) - Universidade Federal do ABC, Programa de Pós-Graduação em Ciências Humanas e Sociais, São Bernardo do Campo.
- UNESCO (2023). *Orientações globais sobre a IA generativa na educação e na pesquisa*. Paris: UNESCO. Disponível em: <<https://unesdoc.unesco.org/ark:/48223/pf0000390241>>. Acesso em: 2 jul. 2025.
- WOOLDRIDGE, Michael (2020). *A Brief History of Artificial Intelligence: where it is, where we are and where we are going*. Nova Iorque: Flatiron Books.

Declaração de Uso de Inteligência Artificial Generativa

O autor do artigo acima identificado declara que ao longo da elaboração deste manuscrito foi utilizada a seguinte ferramenta de inteligência artificial: Claude, com a finalidade de: revisão de linguagem — aprimoramento do estilo textual e da clareza; revisão estilística; verificação de citações; e identificação de padrões. Após o uso da ferramenta, o autor revisou e editou criticamente todo o conteúdo, assumindo total e integral responsabilidade pelo manuscrito publicado, inclusive por eventuais imprecisões ou incorreções.