

Autômatos Digitais de Imagens Médias, Máquinas de Aprendizado Profundo ou Tradutor Parafrástico: Inteligência Artificial entre reflexão e prática estética

Patricia dos Santos (USP)

Resumo

Propõe-se um breve panorama de como a reflexão estética e a prática artística de dois artistas e pensadores nos últimos anos auxiliam na construção de instrumentos tradutórios a respeito dos discursos em voga sobre a natureza e funções da Inteligência Artificial. Partindo das discussões de Ted Chiang e Hito Steyerl, o trabalho questiona o discurso de inovação e criatividade que acompanha a ascensão dos grandes modelos de linguagem e de imagem, sugerindo que tais sistemas funcionam menos como instâncias criadoras e mais como autômatos de repetição estatística: máquinas de aprendizado profundo, mas de compreensão superficial. O texto discute, ainda, a infraestrutura material e econômica dessas tecnologias (seus custos energéticos, a exploração laboral e a lógica especulativa que sustenta sua expansão), evidenciando o paradoxo entre a promessa de imaterialidade e o peso físico do trabalho e da energia que as mantêm.

Palavras-chave: Inteligência Artificial; Estética; Automatismo; Precarização do trabalho.

O avanço das tecnologias de inteligência artificial (IA) nas últimas décadas consolidou uma nova gramática técnica e simbólica do mundo digital. Entre as múltiplas promessas que cercam essas ferramentas (da automação criativa à simulação cognitiva), poucos discursos têm sido tão persistentes quanto o da “inteligência” que cria, interpreta e inova. Esta apresentação busca tensionar tal narrativa, buscando por definições que alguns artistas pensadores, ou seja, profissionais que criam arte, mas também se dedicam a pensar essa criação (e não apenas a sua criação, mas o contexto onde estão inseridos e o trabalho de outros artistas), tem nos dado. Acreditamos que esse prisma pode ser fértil para discutir, por meio de profissionais que trabalham essencialmente com criação, o que é,

se é que existe, e qual a natureza de um potencial criativo na IA, como muitos têm defendido. Imaginamos que, acompanhando as reflexões destes dois artistas de maior destaque, a saber, o escritor Ted Chiang e a artista visual Hito Steyerl, ao questionarmos o que é IA, na intenção de brevemente ensaiarmos algumas definições, consigamos então compreender o que ela é e não é, o que ela pode e não pode oferecer, como ela funciona, e quais são as consequências de maior destaque, num sentido mais abrangente, desse seu funcionamento numa implicação coletiva e social.

A investigação que aqui se propõe é, portanto, dupla: de um lado, analisa-se a estrutura técnica e epistemológica das IAs generativas como dispositivos de compressão e interpolação; de outro, discute-se a dimensão ética e política dessa automatização estética, considerando seus efeitos sobre o trabalho humano, o meio ambiente e a circulação de informação. Aqui buscaremos argumentar que as inteligências artificiais generativas, em especial os grandes modelos de linguagem e imagem, não são dispositivos dotados de criatividade ou compreensão, mas sistemas de compressão com perdas: algoritmos capazes de reconfigurar, com eficiência estatística, uma parcela da cultura humana convertida em dados. Se as inteligências artificiais prometem criar, compreender e decidir, cabe indagar que tipo de mundo elas realmente constroem e, sobretudo, quem se beneficia dessa construção.

A Compressão com Perdas como Paradigma: Ted Chiang e o JPEG Borrado da Web

No debate contemporâneo sobre a natureza e as capacidades dos modelos de linguagem de grande porte, como o ChatGPT, enquanto ferramentas de criação artística, o ensaio do escritor e cientista da computação Ted Chiang, "ChatGPT Is a Blurry JPEG of the Web" ("ChatGPT é um JPEG borrado da Web"), ganhou notoriedade. Isso pois ele nos oferece uma analogia bastante lúcida e reveladora: ele propõe que entendamos esses

sistemas não como mentes emergentes, mas como algoritmos sofisticados de compressão de dados com perdas¹, aplicados ao vasto *corpus* de texto da internet.

Para ilustrar seu ponto, Chiang recorre ao caso de um incidente ocorrido com uma fotocopadora *Xerox*, que utilizava um formato de compressão com perdas (JBIG2) para imagens em preto e branco. O algoritmo identificava regiões similares no documento e, para economizar espaço, armazenava apenas uma versão, reutilizando-a onde julgava adequado. Em um caso específico, ao copiar uma planta baixa com as áreas de três cômodos diferentes (14,13, 21,11 e 17,42 m²), a máquina considerou os números suficientemente parecidos e imprimiu "14,13" em todos os cômodos. O problema crucial, aponta Chiang, não foi tanto o uso da compressão com perdas em si, mas o fato de o artefato gerado, ou seja, o número incorreto, ser legível e plausível, mascarando sua imprecisão sob uma aparência de integridade.

É praticamente isso, argumenta ele, que um modelo como o ChatGPT faz, ao procurar comprimir todo o texto que ele consegue acessar na web, numa espécie de modelo de compressão com perdas (uma vez ser ainda impossível, em termos estruturalmente técnicos, acessar e lidar com toda a informação disponível na internet). Assim, ele identifica "regularidades estatísticas" no texto e as armazena de forma a permitir uma reconstrução aproximada posterior. Ele é, em essência, "um JPEG desfocado de todo o texto da Web". Ele retém a essência informativa, mas não as sequências exatas de bits (ou palavras). Portanto, a excelência do modelo em gerar texto gramatical e fluente ofusca o "desfoque", tornando a aproximação geralmente aceitável.

¹ Aqui talvez valha a pena explicitarmos a inspiração da analogia do autor, que explica: "A compactação de um arquivo requer duas etapas: primeiro, a codificação, durante a qual o arquivo é convertido para um formato mais compacto, e depois a decodificação, em que o processo é revertido. Se o arquivo restaurado for idêntico ao original, o processo de compactação é descrito como sem perdas: nenhuma informação foi descartada. Por outro lado, se o arquivo restaurado for apenas uma aproximação do original, a compactação é descrita como com perdas: algumas informações foram descartadas e agora são irrecuperáveis. A compactação sem perdas é o que normalmente é usado para arquivos de texto e programas de computador, porque esses são domínios nos quais até mesmo um único caractere incorreto tem o potencial de ser desastroso. A compactação com perdas é frequentemente usada para fotos, áudio e vídeo em situações nas quais a precisão absoluta não é essencial. Na maioria das vezes, não percebemos se uma imagem, música ou filme não é reproduzido perfeitamente. A perda de fidelidade se torna mais perceptível apenas quando os arquivos são compactados com muita força. Nesses casos, notamos o que é conhecido como artefatos de compressão: a imprecisão das menores imagens JPEG e MPEG, ou o som metálico de MP3s de baixa taxa de bits"(tradução nossa).

Dessa analogia central, o autor nos explicita, por exemplo, que as chamadas "alucinações" ou invencionices dos modelos (respostas factualmente incorretas, mas usualmente linguisticamente coerentes) são resultados naturais dessa lógica de compressão. Se grande parte do texto original a se orientar for descartado num processo de "compactação" estatística, como ocorre no modelo com perdas, é inevitável que partes da "reconstrução" sejam fabricadas. Assim como a fotocopadora interpolou um número onde não deveria, modelos como o ChatGPT interpolam texto no "espaço lexical" entre conceitos, criando conteúdo novo e, por vezes, enganosamente verossímil. Eles sucedem numa espécie de aproximação estética das palavras num conjunto geral, mas falham no entendimento real, profundo e conceitual de suas criações.

No ensaio de Ted Chiang, a analogia entre os grandes modelos de linguagem e um arquivo JPEG comprimido serve para evidenciar que o funcionamento dessas máquinas não se dá pela apreensão significativa dos mais variados saberes do mundo, mas por meio da redução estatística e sintática do conhecimento humano disponível online. Tal como um arquivo comprimido que preserva apenas a aparência geral da imagem original, sacrificando nuances e detalhes irrecuperáveis, os sistemas de IA generativa processam e reproduzem fragmentos de linguagem de modo a simular coerência, mas sem qualquer consciência do conteúdo semântico que veiculam. Assim, se o conhecimento humano pode ser entendido como um campo heterogêneo de experiências, interpretações e, até mesmo, contradições e ambivalências, a operação das IA's sobre esse campo implica uma planificação estética e epistemológica, produzindo o que Chiang chama, metaforicamente, de uma versão borrada do estoque de um certo saber coletivo, disponibilizado nos arquivos digitais que alimentam esses programas; ou seja, uma imagem de baixa resolução do repositório de um tanto do pensamento humano.

Esse raciocínio ganha força quando relacionado às recentes conclusões da própria OpenAI, que reconhece que as chamadas "*alucinações*" (aquelas suas respostas factualmente incorretas, muitas vezes absurdas e ilógicas, mas linguisticamente bastante convincentes), são matematicamente inevitáveis². Tal constatação confirma, em termos

² "Assim como os alunos que enfrentam questões difíceis em provas, os grandes modelos de linguagem às vezes conjecturam quando estão incertos, produzindo afirmações plausíveis, porém incorretas, em vez de admitir a incerteza; tais 'alucinações' persistem mesmo em sistemas de última geração e minam a confiança", escreveram os pesquisadores da OpenAI Adam Tauman Kalai, Edwin Zhang e Ofir Nachum,

técnicos, a intuição literária e filosófica de Chiang: a estrutura probabilística dos modelos de linguagem, orientada para a verossimilhança e não para a factualidade, os condena a uma forma permanente de erro, inerente ao seu modo de operar. Ao “conjecturar” diante da incerteza, no lugar de assumir seu desconhecimento, o modelo apenas reafirma a natureza de sua compressão: ele não interpreta, apenas aproxima, e o faz de modo a privilegiar a “aparência de sentido” em detrimento de qualquer correspondência com o real ou de assunção de ignorância. Nesse sentido, o que se chama de alucinação é apenas o sintoma mais visível de um mecanismo de cálculo que não pode senão reproduzir a mediocridade dos dados que o alimentam, transformando o conhecimento em um contínuo borramento.

Logo, Chiang contrapõe essa sua ideia dos atuais modelos de IA como JPEG’s borrados da WEB, tão bem expressa por meio de sua analogia ao modelo de compactação com perdas, à visão de que a compressão eficiente equivale à compreensão genuína, já que o modelo produz uma aproximação superficial baseada em padrões estatísticos, e não numa compreensão conceitual. A impressão de que ele “entende” surge precisamente porque ele parafraseia o material da web em vez de citá-lo literalmente. Essa reformulação, que em um estudante humano é sinal de aprendizado, aqui é um subproduto da compressão com perdas, criando a ilusão de que o ChatGPT “entende o material”. Por isso, seu potencial está mais em reformular conteúdos de forma fluida e inteligível, criando a sensação de compreensão, ainda que baseada em aproximações. O desafio, logo, não é só reconhecer a utilidade desses sistemas, mas também compreender os limites dessa tal compressão: o que se ganha em acessibilidade e reformulação é contrabalançado por uma perda, ou seja, pelo risco de perder a precisão factual.

A ilusão de inteligência destes programas emerge, portanto, da forma com que suas respostas são apresentadas. Um texto gramaticalmente coerente cria a aparência de

juntamente com Santosh S. Vempala, da Georgia Tech num estudo recente, publicado em 4 de setembro. Nele, eles argumentam que os modelos de linguagem alucinam porque os procedimentos de treinamento e avaliação recompensam a suposição em vez do reconhecimento da incerteza, e analisam as causas estatísticas das alucinações no processo de treinamento moderno. As alucinações se originam simplesmente de erros na classificação binária. Se afirmações incorretas não podem ser distinguidas de fatos, então as alucinações em modelos de linguagem pré-treinados surgirão por meio de pressões estatísticas naturais. Eles argumentam, então, que as alucinações persistem devido à forma como a maioria das avaliações em testes desses modelos de linguagem é classificada: os modelos de linguagem são otimizados para serem bons candidatos a testes, e suposições quando incertas melhoram o desempenho nos testes. Para conferir a publicação do referido estudo, acessar: <https://arxiv.org/abs/2509.04664>

raciocínio, assim como uma imagem nítida cria a aparência de realismo. Entretanto, o que se preserva não é tanto o conteúdo do conhecimento, mas o seu efeito de legibilidade. Essa lógica aproxima os modelos de IA de uma espécie de processo de tradução automática da cultura humana hegemônica: eles não interpretam, mas reformulam; não compreendem, mas interpolam. Trata-se de um tipo de tradução parafrástica, na qual o sentido é substituído por um simulacro de familiaridade. Nesse sentido, as IAs generativas não são entidades cognitivas, mas autômatos tradutores de uma imensa base de dados. Sua suposta criatividade é o reflexo de uma compressão cultural em larga escala, na qual o “novo” é apenas uma recombinação estatística do já existente. A metáfora da compressão, assim, expõe não apenas a limitação epistemológica dessas tecnologias, mas também o modo como o capitalismo informacional transforma o conhecimento em matéria-prima: compactando-o, armazenando-o e redistribuindo-o como valor econômico.

Daí que deve-se entender que a IA trabalha por meio da seleção estatística de um determinado estoque de dados e que, nesse sentido, ela, ainda, é incapaz de praticar uma ação, que de maneira alguma é específica a humanidade, mas é sim específica aos seres vivos, que é a da capacidade criativa real e total. Dessa forma, Ted Chiang nos leva a uma compreensão da IA, enquanto ferramenta de escrita, como uma espécie de *Tradutor Parafrástico*. Por isso que, para o escritor, chamar essas ferramentas de Inteligência Artificial, ou numa linguagem mais técnica, de Máquinas de Aprendizado Profundo, é, no mínimo incorreto, pelo menos quando se tem uma concepção mais valorativa do conceito de aprendizado profundo³. Até porque um aprendizado profundo que se leve a sério deve ser capaz, no mínimo, de reconhecer erros, assumir limitações e desconhecimentos, e então aprender com isso - algo que esses modelos não tem feito.

³ Não perdemos de vista a oportunidade exemplar que aí nos é oferecida de discutirmos o que são certos conceitos como “inteligência”, “conhecimento” e “aprendizado”. Porém, tendo em mente a brevidade deste texto e que nenhum dos autores principais dos quais pretendemos aqui apresentar entram naquilo que entendem mais profundamente como definições destes conceitos, optamos por não abarcar tal tarefa.

Da imagem média à imagem automatizada: Hito Steyerl e o redundante imaginário algorítmico

Em junho de 2023, a artista visual Hito Steyerl publicou um notável ensaio chamado *Mean Images*, no qual procura descrever a emergência de uma visualidade típica da era digital, marcada pela homogeneização estética e pela padronização algorítmica, isso por meio da produção visual operada através de Inteligência Artificial. Ela nomeia a típica e massivamente produzida imagem gerada por modelos de aprendizado profundo como “imagens médias”: elas são aquelas que resultam de uma estatística do olhar; são produtos de tecnologias que analisam e recompõem o mundo a partir da média de milhares de exemplos; imagens “otimizadas” para o consumo, calculadas para se aproximar do que é mais provável, mais reconhecível, e supostamente mais desejável.

O ensaio se estrutura num jogo de palavras que nos permite entender esse título e o próprio conceito que ela leva ao longo do texto tanto como “imagens médias”, “imagens medianas/mediócras” e “imagens más”. Nesse texto ela propõe uma chave de leitura que não é apenas conceitual, mas também política: as imagens geradas por IA não são singulares nem originais, justamente porque elas derivam de cálculos de probabilidade que extraem regularidades estatísticas de vastos conjuntos de dados, assim como explicitado por Ted Chiang com as IAs que lidam com palavras e textos. Assim, essas imagens são “médias” no sentido matemático e, ao mesmo tempo, “más” no sentido qualitativo: porque são medianas, mediócras, banais, frequentemente insatisfatórias, e reveladoras de uma mediocridade estrutural da WEB onde essas ferramentas extraem seus estoques, estoques estes que funcionam como coordenadas do seu funcionamento.

Esse duplo caráter desses modelos (estatístico e valorativo), abre para a autora um campo fértil de discussão. Porque, de um lado, as imagens médias revelam o modo de funcionamento da IA como dispositivo de repetição e cálculo, que não cria a partir do nada, mas recicla o que já foi ou é dado. Do outro lado, elas mostram a face política de uma cultura que tende à homogeneização, ao apagamento das singularidades, e ao nivelamento do sensível. Nesse movimento, a reflexão estética proposta por Hito se torna inseparável de uma análise crítica da sociedade contemporânea, marcada pela financeirização, pela lógica de *big data* e pela precarização do trabalho. Ou seja, para a

artista, a IA, como tem se desenvolvido até o presente momento, pode ser entendida como uma espécie de *Autômato Digital de Imagens Médias e Más*.

Fica claro que o núcleo conceitual deste ensaio está no jogo de palavras entre “média” e “má”, enquanto adjetivação da imagem, desse produto visual gerado por IA: o raciocínio que a autora desenha, dentro desse universo que procura analisar, salienta que em termos estatísticos, a média é um cálculo de tendência central, que busca condensar um conjunto de dados em um valor representativo. Em termos estéticos e sociais, a média remete à mediocridade, ao comum, ao banal. Essa coincidência não é casual: a imagem gerada por IA tende a ser justamente e exatamente isso — uma representação que não se destaca, que resulta da soma e do balanceamento de inúmeras imagens pré-existentes, tornando-se então um “compromisso” estatístico.

Dessa forma, Hito demonstra que o aprendizado de máquina de IA's opera por aproximação, e não por invenção. Esse aprendizado pondera recorrências, combina probabilidades, e ajusta padrões. O que emerge daí não é o singular, mas o provável; e o provável, nesse contexto, é justamente o que já ocorreu muitas vezes antes, e que não é sequer necessariamente o correto⁴. Então a originalidade, que muitas vezes é entendida como ruptura ou desvio, dificilmente encontra espaço num regime que se orienta pelo cálculo do mais frequente. Assim, a imagem de IA tende à redundância e à previsibilidade.

Em seu ensaio, portanto, Hito Steyerl descreve a emergência de uma visualidade típica da era digital, marcada pela homogeneização estética e pela padronização algorítmica. As imagens médias são aquelas que resultam de uma estatística do olhar, ou seja, produtos de tecnologias que analisam e recompõem o mundo a partir da média de milhares de exemplos. São imagens “otimizadas” para o consumo, calculadas para se aproximar do que é mais provável, mais reconhecível, mais desejável.

⁴ Como a questão das alucinações, já comentadas previamente neste texto provam. Para um breve comentário sobre a questão das alucinações de IA, num viés informativo e jornalístico, recomendamos a reportagem “OpenAI admits AI hallucinations are mathematically inevitable, not just engineering flaws” (“OpenAI admite que as alucinações da IA são matematicamente inevitáveis, não apenas falhas de engenharia”), de autoria de Gyana Swain, publicada na revista online ComputerWorld, disponível em: <https://www.computerworld.com/article/4059383/openai-admits-ai-hallucinations-are-mathematically-inevitable-not-just-engineering-flaws.html>

Quando transposta para o contexto das inteligências artificiais generativas, essa noção se torna ainda mais contundente. Os sistemas de geração de imagem, como DALL·E ou Midjourney, constroem o visível a partir de uma imensa base de dados pré-existente. Cada resultado é, no fundo, uma média visual: uma aproximação estatística do imaginário coletivo digital. Isso porque a IA não “vê o mundo” para então criar suas representações, mas o reconstrói com base em imagens já vistas, reiterando os padrões dominantes de representação. Assim, o automatismo da produção algorítmica tende a reforçar o que já é familiar, consolidando um regime estético da redundância.

A estética das imagens médias é também a estética da repetição e da convergência: quanto mais as IAs produzem imagens, mais essas imagens alimentam seus próprios bancos de dados, retroalimentando o mesmo repertório formal e semântico. O resultado é um círculo vicioso de homogeneização cultural, no qual a diferença se dissolve na previsibilidade e o singular se dilui no genérico. Steyerl identifica nesse processo uma forma de automatismo digital, que desloca o artista para a posição de operador de parâmetros, enquanto o algoritmo assume a autoria formal.

Contudo, essa automatização da forma não se traduz em uma verdadeira autonomia da máquina. O algoritmo continua dependente daquilo que já foi produzido por sujeitos humanos, cuja criatividade é reabsorvida e redistribuída como dado. A IA, nesse sentido, não cria novas imagens do mundo, mas replica estatisticamente o mundo das imagens. Suas produções são, como propõe Steyerl, imagens médias e medíocres; médias no sentido técnico, e medíocres no sentido político: mediadoras de um imaginário coletivo cada vez mais capturado por lógicas de predição e consumo.

Economia política e infraestruturas da IA: trabalho, energia e especulação

Para além do interessante e produtivo conceito de “Imagens Médias”, o ensaio de Hito também traz uma contribuição notável no campo das reflexões contemporâneas críticas sobre IA, num prisma mais materialista e objetivo, principalmente quando dedica-se a comentar sobre a questão do trabalho humano que sustenta o desenvolvimento destas tecnologias. Por trás da aparente imaterialidade das inteligências artificiais, há uma

infraestrutura massiva de atividade humana, exploração territorial e consumo energético. O discurso tecnocrático que apresenta a IA como uma instância autônoma e autoaprendente depende de uma enorme cadeia invisível de trabalhadores: programadores, moderadores de conteúdo, etiquetadores de dados e operadores de centros de treinamento, muitos deles localizados em países do Sul Global. Hito Steyerl enfatiza que o funcionamento das IAs só é possível porque há uma força de trabalho precarizada que treina, corrige e alimenta os modelos, ainda que a narrativa dominante tente apagar essa dimensão humana do processo. A suposta “inteligência” das máquinas depende, paradoxalmente, de um imenso esforço cognitivo humano, realizado em condições laborais degradantes e muitas vezes traumáticas. Assim, o mito da IA autossuficiente se revela uma versão digital do velho princípio colonial: a acumulação de valor em um polo implica a exaustão material e psíquica em outro.

A artista defende que essa ocultação do trabalho humano é parte de uma lógica colonial contemporânea, na qual o saber técnico e a infraestrutura produtiva são geograficamente separados. A revolucionária eficiência do “aprendizado profundo” repousa sobre o trabalho invisível de milhares de pessoas mal remuneradas, responsáveis por transformar o caos informacional da internet em dados úteis. Como observa a autora, o digital não elimina o corpo, mas o redistribui, externalizando suas demandas físicas e energéticas para zonas de invisibilidade social e geográfica. É justamente nesse ponto que a IA contemporânea se revela como um dispositivo intrinsecamente político: ela depende de vastos sistemas de exploração de recursos (humanos, minerais e ambientais) que são sistematicamente apagados do discurso tecnológico hegemônico.

Paralelamente, a sustentação física da IA exige um custo energético e ecológico crescente. Os grandes data centers consomem quantidades exorbitantes de eletricidade e água para refrigeração, intensificando o impacto ambiental dessas tecnologias. A promessa de uma inteligência “etérea” se revela, assim, ancorada em uma infraestrutura material de alta densidade, cuja manutenção é incompatível com qualquer ideal de sustentabilidade. O poder computacional necessário para treinar modelos de IA generativos⁵, pode exigir uma quantidade impressionante de eletricidade, o que leva ao

⁵ Um outro destacado problema relacionado aos danos ambientais causados direta e indiretamente por modelos de Inteligência Artificial é o da dificuldade de calculá-los com precisão. Não apenas porque o poder computacional necessário para treinar modelos de IA generativos geralmente têm bilhões de

aumento das emissões de dióxido de carbono e pressões na rede elétrica. Além da demanda por eletricidade, a grande quantidade de água necessária para resfriar o hardware usado para treinar, implantar e ajustar modelos de IA generativa, pode sobrecarregar o abastecimento de água municipal e afetar os ecossistemas locais⁶. O número crescente de aplicações de IA generativa também impulsionou a demanda por hardware de computação de alto desempenho, adicionando impactos ambientais indiretos de sua fabricação e transporte.

Embora nem toda computação de data center envolva IA generativa, a tecnologia tem sido um dos principais impulsionadores do aumento da demanda de energia. Como destacado pelo cientista Noman Bashir,⁷ a demanda por novos data centers “não pode ser atendida de forma sustentável”, pois o ritmo em que as empresas estão construindo novos data centers significa que “a maior parte da eletricidade para abastecê-los precisa vir de usinas de energia movidas a combustíveis fósseis”⁸. Embora todos os modelos de

parâmetros, mas principalmente porque as empresas responsáveis por eles têm sido especialmente cautelosas no oferecimento público destas informações. Essa falta de transparência, porém (que só pode parecer proposital, na intenção de diminuir ao máximo as críticas ao desenvolvimento dessas ferramentas), não tem impedido pesquisadores e jornalistas de investigar e calcular a verdadeira dimensão dos impactos ambientais da IA que se alastram pelo mundo. Cientistas estimam que a demanda de energia dos data centers na América do Norte aumentou de 2.688 megawatts no final de 2022 para 5.341 megawatts no final de 2023, em parte impulsionada pelas demandas da IA generativa. Globalmente, o consumo de eletricidade dos data centers aumentou para 460 terawatts-hora em 2022. Isso tornaria os data centers o 11º maior consumidor de eletricidade do mundo, entre os países Arábia Saudita (371 terawatts-hora) e França (463 terawatts-hora), de acordo com a Organização para a Cooperação e Desenvolvimento Econômico (OCDE). Até 2026, espera-se que o consumo de eletricidade dos data centers se aproxime de 1.050 terawatts-hora (o que colocaria os data centers em quinto lugar na lista global, entre Japão e Rússia).

⁶ Embora a demanda de eletricidade dos data centers possa estar recebendo mais atenção na literatura de pesquisa, a quantidade de água consumida por essas instalações também tem impactos ambientais. Água gelada é usada para resfriar um data center, absorvendo o calor dos equipamentos de computação. Estima-se que, para cada quilowatt-hora de energia consumido por um data center, seriam necessários dois litros de água para resfriamento, diz Bashir. Os dados referentes aos impactos ambientais do desenvolvimento e uso de Inteligência Artificial e o relato do pesquisador Norman Bashir foram extraídos da reportagem “Explained: Generative AI’s environmental impact”, publicada no site de notícias do Massachusetts Institute of Technology, que pode ser acessada em: [Explained: Generative AI’s environmental impact | MIT News | Massachusetts Institute of Technology](#)

⁷ Noman Bashir é pesquisador de Computação e Impacto Climático no Consórcio de Clima e Sustentabilidade do MIT (MCSC) e pós-doutorado no Laboratório de Ciência da Computação e Inteligência Artificial (CSAIL).

⁸ Tomemos, por exemplo, mais um dado retirado da reportagem da qual extraímos o importante relato de Bashir: a energia necessária para treinar e implementar um modelo como o GPT-3 da OpenAI é difícil de determinar, mas em um artigo de pesquisa de 2021, cientistas do Google e da Universidade da Califórnia em Berkeley estimaram que apenas o processo de treinamento consumiu 1.287 megawatts-hora de eletricidade (o suficiente para abastecer cerca de 120 residências médias nos EUA por um ano), gerando cerca de 552 toneladas de dióxido de carbono.

aprendizado de máquina devam ser treinados, um problema exclusivo da IA generativa são as rápidas flutuações no uso de energia que ocorrem em diferentes fases do processo de treinamento. Os operadores da rede elétrica precisam ter uma maneira de absorver essas flutuações para proteger a rede e geralmente empregam geradores a diesel para essa tarefa⁹.

Ainda que não abordado por Hito em seu ensaio, não pode-se perder de vista a obviedade com a qual, a esse quadro, soma-se o problema da dimensão especulativa do investimento em IA. A corrida pelo desenvolvimento de modelos cada vez maiores cria uma bolha de expectativa¹⁰, sustentada por promessas de rentabilidade futura e pela mitologia do progresso técnico. Trata-se de uma economia da fé, de uma confiança no potencial quase messiânico da tecnologia para resolver crises, crises estas que frequentemente deixa-se de perceber que a mesma ajuda a agravar, como no caso da crise climática.

Assim, a retórica apocalíptica sobre o “poder destrutivo” da IA, propagada por muitos de seus criadores e maiores defensores, serve muitas vezes para inflar seu valor de mercado, convertendo o medo em capital. Ao revesti-la de um enorme potencial destruidor, discurso ainda que sempre acompanhado de garantias de que os responsáveis por sua criação saberão como impedi-lo, inflaciona-se suas reais capacidades, para que assim continue-se obtendo investimentos vultosos para criação de versões pretensamente

⁹ À este cenário, acrescenta-se mais um dado: depois que um modelo de IA generativa é treinado, as demandas de energia não desaparecem. Cada vez que um modelo é utilizado, talvez por um indivíduo que solicita ao ChatGPT um resumo de um e-mail, o hardware de computação que realiza essas operações consome energia. Pesquisadores estimam que uma consulta ao ChatGPT consuma cerca de cinco vezes mais energia do que uma simples pesquisa na web. “Mas um usuário comum não pensa muito nisso”, comenta Bashir. “A facilidade de uso das interfaces de IA generativa e a falta de informações sobre os impactos ambientais das minhas ações significam que, como usuário, não tenho muito incentivo para reduzir o uso de IA generativa.”

¹⁰ Progressivamente, aumenta-se o ritmo de publicações de artigos no noticiário internacional a respeito de uma possível e provável bolha da IA, e seu consequente estouro. Tome-se, por exemplo, o fato de que dois dias antes do envio deste texto para publicação, em 18 de outubro de 2025, foi publicado no site da CNN uma reportagem intitulada “Por que este analista diz que a bolha da IA é 17 vezes maior que a bolha da Internet” (“Why this analyst says the AI bubble is 17 times bigger than the dot-com bust”). Nela, ainda que saliente-se que só naquela semana, o Financial Times havia noticiado que 10 startups de IA — nenhuma delas com qualquer margem de lucro sequer — haviam ganho quase US\$ 1 trilhão em valor de mercado nos últimos 12 meses, o foco se dava nos alertas que o pesquisador e sócio da empresa britânica MacroStrategy Partnership Julien Garran, dava num relatório recentemente publicado, no qual alertava que encontrávamos na “maior e mais perigosa bolha que o mundo já viu”, concluindo que há uma “má alocação de capital nos EUA” que torna o frenesi atual em volta do mercado de IA 17 vezes maior que a bolha da Internet e quatro vezes maior que a bolha imobiliária de 2008.

mais evoluídas¹¹. Com anos de acúmulo de perdas e prejuízos financeiros, e pesquisas que parecem avolumar-se nas indicações de que a satisfação de seus contratantes tem sido cada vez menor, o discurso que inflaciona as capacidades reais das IA's torna-se progressivamente essencial para manter o interesse de investidores e subsidiar seu desenvolvimento¹². O resultado é uma forma de economia estética e política na qual a imagem da inteligência e seu discurso de excelência e ineditismo técnico vale mais do que sua efetiva capacidade de pensar e produzir eficiência.

Direitos autorais, retroalimentação e o problema do estoque

Como exemplificado pelas ideias de Chiang e Steyerl, a base de funcionamento das inteligências artificiais contemporâneas é o *estoque*, um arquivo gigantesco de dados textuais, visuais e sonoros que serve de alimento a seus algoritmos. Toda a operação de IA depende de uma extração contínua e massiva de informações disponíveis na internet, muitas vezes retiradas de modo indiscriminado, incluindo obras protegidas por direitos autorais e produções culturais que procuraram se destacar como expressão de experiências humanas singulares. Essa dependência estrutural de um repositório prévio

¹¹ O desenvolvimento progressivo de versões atualizadas destes modelos de aprendizado profundo também carregam seu custo ecológico: com a IA tradicional, o consumo de energia é dividido de forma bastante uniforme entre processamento de dados, treinamento de modelos e inferência, que é o processo de usar um modelo treinado para fazer previsões com base em novos dados. No entanto, como lembra o cientista Norman Bashir, espera-se que as demandas de eletricidade da inferência de IA generativa acabem dominando, visto que “esses modelos estão se tornando onipresentes em muitas aplicações, e a eletricidade necessária para a inferência aumentará à medida que versões futuras dos modelos se tornarem maiores e mais complexas”. Além disso, os modelos de IA generativa têm uma vida útil particularmente curta, impulsionada pela crescente demanda por novas aplicações de IA. “As empresas lançam novos modelos a cada poucas semanas, então a energia usada para treinar versões anteriores é desperdiçada”, acrescenta Bashir. Novos modelos geralmente consomem mais energia para treinamento, pois geralmente têm mais parâmetros do que seus antecessores.

¹² Como comentado pelo economista Michael Roberts, em texto intitulado “The AI bubble and the US economy” (“A bolha da IA e a economia dos EUA”), publicado em seu blog The next recession (disponível em: <https://thenextrecession.wordpress.com/2025/10/14/the-ai-bubble-and-the-us-economy/>). até agora, há poucos sinais de que o investimento em IA esteja proporcionando grandes aumentos de produtividade. Porém, ironicamente, o enorme investimento em centros de dados e em infraestrutura de IA está sustentando a economia dos EUA nesse meio tempo. Quase 40% do crescimento real do PIB dos EUA no último trimestre foi impulsionado pela despesa de capital em tecnologia e a maior parte dele ocorreu em investimentos relacionados à IA. A infraestrutura de IA recebeu investimentos de US\$ 400 bilhões desde 2022.

evidencia uma contradição central: as máquinas são apresentadas como autônomas e criativas, mas seu “pensamento” é apenas uma recombinação estatística do que já existe. A IA não cria a partir do nada, pois ela rearranja o que está armazenado, comprimindo o conteúdo do mundo num arquivo que, como um JPEG em baixa resolução, degrada a riqueza e a nuance do conhecimento humano.

Essa lógica de estoque também explica a ironia ética e política do atual regime informacional. As mesmas corporações que hoje se apropriam de forma silenciosa e não remunerada do trabalho coletivo da internet são aquelas que defendem uma rede cada vez mais privatizada e monetizada. O ideal utópico de um ciberespaço livre, sustentado por ativistas como Aaron Swartz, foi substituído por um ecossistema corporativo que transforma cada fragmento de cultura em dado e cada dado em propriedade. A IA, nesse contexto, devora o comum sem devolvê-lo, alimentando-se de um patrimônio coletivo de ideias, estilos e linguagens para gerar produtos cuja circulação e lucro permanecem concentrados nas mãos de poucos. Paradoxalmente, o sonho de uma internet democrática, voltada ao livre acesso e à colaboração, renasce dentro das máquinas apenas como simulacro: a IA encena a inteligência coletiva, mas a privatiza.

Com o tempo, esse processo de alimentação do estoque se volta sobre si mesmo. À medida que os modelos de IA começam a ser treinados com dados gerados por outras IA's, a rede entra em um estado de retroalimentação entrópica: o estoque passa a reproduzir o próprio estoque¹³. Esse circuito fechado transforma o espaço digital em um

¹³ Como mais uma prova à teoria popular de que a Internet está morta, surpreende a velocidade com a qual se alastra os materiais produzidos por IA em todas as frentes da WEB. Como destacado no artigo “A IA está se matando silenciosamente – e matando a Internet?” (“Is AI quietly killing itself – and the Internet?”), de autoria de Tor Constantino, publicado no site da Forbes australiana, cerca de 57% de todo o texto da web foi gerado por IA ou traduzido por meio de um algoritmo de IA, de acordo com um estudo conduzido por uma equipe de pesquisadores da Amazon Web Services, intitulado “A Shocking Amount of the Web is Machine Translated: Insights from Multi-Way Parallelism” (disponível em: <https://arxiv.org/pdf/2401.05749>). Esse dado é preocupante, dentre muitos motivos, pelo fato de que quando IA's são alimentadas por IA's, suas respostas começam a rapidamente degradar. Muitos estudos sobre essa questão mostram que, depois de alguns ciclos de treinamento com esse tipo de dado, os modelos passam a cometer erros significativos. Depois, produzem informações sem sentido algum. Inicialmente, os dados não compreendidos pela IA original são subrepresentados na IA que usa suas informações, até que os erros aumentam e se reproduzem, inutilizando os modelos, atingindo, por fim, um estado total de colapso. Logo, se os dados gerados por humanos na internet estão sendo rapidamente encobertos por conteúdo gerado por IA, tendo em vista as conclusões dos estudos referidos, é possível supor um caminho pelo qual a IA vá progressivamente se autodestraindo e, por consequência, a internet. (conferir: <https://www.forbes.com.au/news/innovation/is-ai-quietly-killing-itself-and-the-internet/>)

ambiente de autossimulação, onde cada vez mais a novidade se dissolve na redundância e a diferença é substituída por repetições probabilísticas. O erro e o clichê tornam-se inevitáveis, pois o sistema aprende consigo mesmo, como uma fotocopadora reproduzindo cópias de cópias. O resultado é um universo simbólico cada vez mais mediano — culturalmente empobrecido e eticamente opaco — em que a produção de sentido é reduzida a uma estatística de frequência. As chamadas “alucinações” dos modelos de linguagem, como demonstram os próprios estudos da OpenAI, não são falhas técnicas, mas consequências matematicamente inevitáveis de um sistema baseado em compressão e recombinação de um estoque finito de informação.

O problema da autoria, portanto, é apenas um sintoma de uma questão mais profunda: a da condição ontológica da criação em tempos de estoque total. Se toda forma de produção depende de um arquivo prévio, o artista, humano ou maquínico, deixa de ser um criador e se torna um operador de um repositório cultural. O gesto criativo é absorvido pelo cálculo, e a autoria dissolve-se no fluxo estatístico das recombinações possíveis. O estoque, por sua vez, precisa ser continuamente expandido, e é nesse ponto que emergem suas consequências materiais e sociais mais graves: a necessidade de novos dados alimenta uma indústria de exploração humana e ambiental, na qual trabalhadores precarizados etiquetam, validam e filtram conteúdos em condições degradantes, enquanto *data centers* consomem energia e água em escala planetária. A IA, portanto, pode também ser definida menos como uma “máquina de pensar” do que um organismo extrativista, dependendo de corpos, territórios e matérias para sustentar a ficção de sua autonomia.

Criação, crítica e uso artístico da IA

Caracterizar a inteligência artificial apenas como destrutiva seria reduzir seu campo de possibilidades críticas. Como toda tecnologia, ela é ambígua: pode servir tanto à reprodução de estruturas de dominação quanto à reflexão sobre elas. A história das mídias mostra que dispositivos inicialmente percebidos como mecânicos ou inexpressivos (como a fotografia ou o cinema), tornaram-se instrumentos de experimentação estética e política. O mesmo pode ocorrer com a IA, desde que ela seja abordada não como substituto da

imaginação humana, mas como espelho de suas condições de produção, como explicitado por Hito Steyerl em seu ensaio. O verdadeiro uso artístico da IA reside em expor seu estoque: em revelar as camadas de extração, os filtros invisíveis, os vieses e as lacunas que estruturam seu funcionamento.

Nesse sentido, há uma continuidade entre as posturas críticas de artistas como Harun Farocki e Jean-Luc Godard e as experiências contemporâneas de criação com IA. Farocki, ao denunciar o caráter funcional das “imagens operativas” (aquelas produzidas por máquinas para outras máquinas)¹⁴, antecipou a crítica à lógica algorítmica que hoje domina a visualidade digital. Godard, por sua vez, viu no vídeo e na imagem eletrônica não apenas uma ameaça em sua direta ligação com o discurso publicitário, mas uma oportunidade de democratização estética, de descentralização da fala e do olhar. Inspirados por esses precedentes, alguns artistas atuais transformam a IA em um campo de desmontagem: desconstruem suas promessas de inteligência, evidenciam suas dependências infraestruturais e reconfiguram suas saídas automatizadas como matéria para uma nova crítica das imagens. Parece-nos que Hito Steyerl é, atualmente, um dos maiores e melhores exemplos desse grupo de artistas.

Mas é preciso reconhecer que todo uso da IA carrega consigo a sombra de seu estoque. Ao operar sobre bancos de dados alimentados por trabalho explorado e energia dispendida, nenhuma criação pode se declarar eticamente neutra. As práticas mais interessantes são justamente aquelas que confrontam esse impasse, que recusam a estetização ingênua da tecnologia e fazem dela um campo de disputa simbólica. Usar a IA criticamente significa politizar seu uso, revelar o que ela oculta: o corpo humano, o custo ecológico, a ideologia da eficiência. O desafio não é produzir belas imagens com a máquina, mas compreender o que ela revela e, principalmente, o que ela apaga sobre o nosso próprio regime de ver, saber e criar.

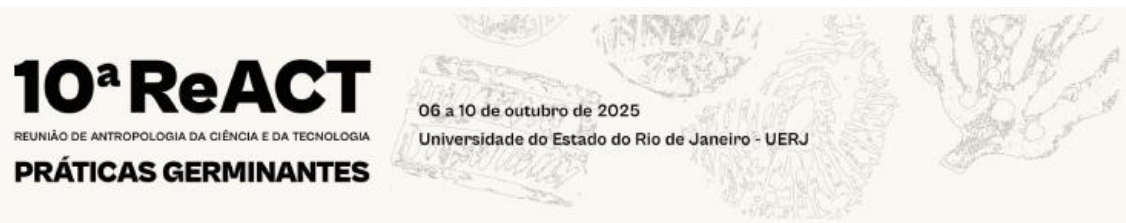
¹⁴ Ampliando a explicação do referido conceito, as “imagens operativas” de Harun Farocki são imagens que não representam um objeto para reflexão, mas sim fazem parte de uma operação ou processo. Cunhado em sua trilogia *Eye/Machine* (2001-2003), o conceito descreve visualizações de dados como as de sistemas de mira de mísseis guiados a laser, inspeção de tubulações de esgoto ou cirurgia robótica, que são criadas para “executar trabalho” e não para serem vistas por humanos. São imagens funcionais, presentes em aplicações militares e civis, que transformam o espectador em mero operador em uma relação mediada pela máquina.

Conclusão: o estoque e a moral da máquina

Como visto, o ponto de partida de nossa apresentação foi a metáfora proposta por Ted Chiang, para quem os grandes modelos de linguagem funcionam como “JPEGs borrados do conhecimento humano”: representações reduzidas que preservam contornos reconhecíveis, mas sacrificam nuance, contexto e profundidade em favor da generalização. Essa analogia nos permitiu compreender que as respostas produzidas por sistemas como o ChatGPT ou o Midjourney não derivam de uma operação intelectual, mas de uma interpolação entre padrões aprendidos. A ideia de inteligência artificial, portanto, encontra-se intrinsecamente vinculada à noção de perda — de singularidade, autoria e sentido —, o que desloca o debate para o campo da estética e da política da representação.

Nesse horizonte, a contribuição de Hito Steyerl tornou-se fundamental. Sua formulação sobre as *imagens médias* descreve a tendência das tecnologias digitais em gerar visualidades padronizadas, mediadas por estatísticas, algoritmos e convenções técnicas que definem a aparência do mundo. Ao transpor esse conceito para o contexto das IAs generativas, é possível sugerir que o que está em jogo não é a invenção de novas formas de ver, mas a automatização de uma estética do consenso; uma visualidade probabilística que reduz a experiência do sensível à previsibilidade do cálculo.

Assim, podemos concluir que as chamadas inteligências artificiais, ao condensarem o conhecimento humano em grandes bases de dados, funcionam como autômatos de estoque, sistemas que armazenam e recombina o mundo sem, de fato, compreendê-lo. Sua inteligência é uma tradução estatística da experiência, e seu erro, as ditas alucinações inevitáveis, as redundâncias e os vieses, é apenas o reflexo lógico de seu modo de operar. Pode-se entender então que a exclusão da moralidade da inteligência maquínica converteu o pensamento em cálculo e a cognição em estatística. Isso pois, ao acreditarem ser possível colocarem as máquinas para processar de maneira ‘neutra’ os dados estocados, a maior parte dos cientistas e desenvolvedores de IA’s afirmam que estas suas criações geram resultados e respostas verdadeiras. No entanto, essas “verdades” são



apenas versões otimizadas daquilo que já circula: repetições de um mundo filtrado e mediado por interesses econômicos e padrões culturais dominantes.

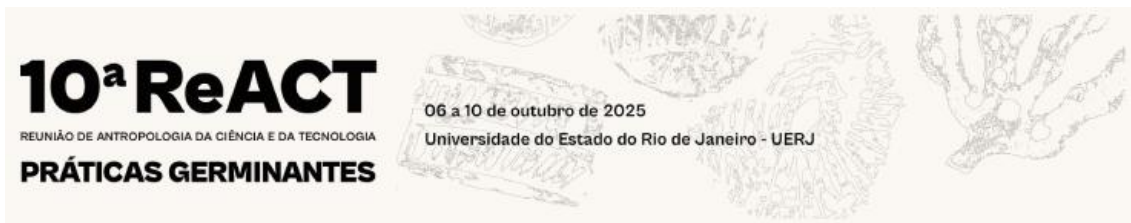
Repensar a IA a partir da noção de estoque é, portanto, um exercício de desmistificação. Significa compreender que sua força e seu perigo provêm do mesmo lugar: da capacidade de acumular e reproduzir informação. O desafio ético e político do presente é interromper a cadeia de reprodução automática — questionar quem alimenta o estoque, com que trabalho e com que custo. O gesto crítico mais radical não é ensinar a máquina a criar, mas ensinar a sociedade a não confundir acúmulo e produtivismo com pensamento.

Bibliografia

BASHIR, Noman. *Explained: Generative AI's environmental impact*. MIT News, 17 jan. 2025. Disponível em: <https://news.mit.edu/2025/explained-generative-ai-environmental-impact-0117>. Acesso em: 18 out. 2025.

CHIANG, Ted. *ChatGPT Is a Blurry JPEG of the Web*. The New Yorker, 9 fev. 2023. Disponível em: <https://www.newyorker.com/tech/annals-of-technology/chatgpt-is-a-blurry-jpeg-of-the-web>. Acesso em: 18 out. 2025.

CONSTANTINO, Tor. *Is AI quietly killing itself — and the Internet?* Forbes Australia, 2024. Disponível em: <https://www.forbes.com.au/news/innovation/is-ai-quietly-killing-itself-and-the-internet/>. Acesso em: 18 out. 2025.



FAROCKI, Harun. *Eye/Machine* [Trilogia]. 2001-2003.

GARRAN, Julien. *Why this analyst says the AI bubble is 17 times bigger than the dot-com bust*. CNN Business, 18 out. 2025. Disponível em: <https://edition.cnn.com/2025/10/18/business/ai-bubble-dot-com/index.html>. Acesso em: 18 out. 2025.

GENARO, E. de. (2015). "Imagens operativas e pós-fotográficas: um estudo a partir de Farocki". *Revista Eco-Pós*, 18(2), 134–150. <https://doi.org/10.29146/eco-pos.v18i2.2345>

KALAI, Adam Tauman et al. *Why Do Language Models Hallucinate? The Statistical Pressures Leading to Model Confabulation*. arXiv, 2025. Disponível em: <https://arxiv.org/abs/2509.04664>. Acesso em: 18 out. 2025.

ROBERTS, Michael. *The AI bubble and the US economy*. The Next Recession Blog, 14 out. 2025. Disponível em: <https://thenextrecession.wordpress.com/2025/10/14/the-ai-bubble-and-the-us-economy/>. Acesso em: 18 out. 2025.

STEYERL, Hito. *Mean Images*. 2023. Disponível em: https://monoskop.org/images/1/19/Steyerl_Hito_2023_Mean_Images.pdf. Acesso em: 18 out. 2025.

SWAIN, Gyana. *OpenAI admits AI hallucinations are mathematically inevitable, not just engineering flaws*. ComputerWorld, 9 set. 2025. Disponível em:

10ª ReACT

REUNIÃO DE ANTROPOLOGIA DA CIÊNCIA E DA TECNOLOGIA

PRÁTICAS GERMINANTES

06 a 10 de outubro de 2025

Universidade do Estado do Rio de Janeiro - UERJ

<https://www.computerworld.com/article/4059383/openai-admits-ai-hallucinations-are-mathematically-inevitable-not-just-engineering-flaws.html>. Acesso em: 18 out. 2025.