

Regulando a Inteligência Artificial: É possível fazer diferente?

Cristiano Galafassi

28-10-2024

Resumo

Este texto aborda a inovação trazida pela Inteligência Artificial (IA) e discute caminhos para a sua regulação. São apresentadas sugestões que contemplam a regulação por design, na origem com base na própria IA e o papel da educação para a regulação voluntária. Ao abordar a educação, se faz necessário questionar os limites da inovação marcadamente tecnológica frente a uma formação humana ampla. Ou seja, como pensar a inovação tecnológica para além do produtivismo. Mas, também, considera o que a educação pode fazer pela tecnologia. Assumindo-se que é necessário regular aspectos da IA, quais as possibilidades para tal? Precisamos da IA? É possível fazer diferente?

Palavras-chave: Regulação da IA, Ética na IA, Educação para a IA, Governança da IA, impactos sociais, econômicos e ambientais da IA.

INTRODUÇÃO

Certamente a IA inova, ela pode trazer várias transformações para a vida individual e coletiva. Pode e está mudando o que entendemos por democracia (Schneier; Sanders, 2023). Esta mudança acontece pelos algoritmos que manipulam ideias e conceitos nas redes sociais, mas também pelo fato de sistemas de IA estarem assumindo parte do papel dos legisladores. Plataformas recebem solicitações diretas dos cidadãos, agrupam as solicitações semelhantes, conferem incompatibilidades com leis já existentes, buscam legislações internacionais similares, geram o texto do projeto de Lei e o encaminham para votação (David *et al.*, 2017) (Salomão, 2022). Como garantir que essas tecnologias não sejam uma ameaça à democracia? Este e outros tantos argumentos são utilizados para justificar a regulação da IA, assegurando que ela sirva ao bem do planeta de forma transparente e responsável.

Este artigo parte da premissa tecnológica de que a inovação da IA não se limita a ser apenas uma ferramenta tecnológica. Ela é um agente dotado de conhecimento, raciocínio, objetivos a serem alcançados e poder de decisão. Um dos desafios é a

difficuldade de explicar como alguns sistemas de IA chegam a uma decisão. Ou seja, a possibilidade de compreender e justificar como ela alcança seus objetivos. Por um lado, a IA possui objetivos explícitos, diretamente programados por seres humanos, que demarcam seu domínio e suas capacidades. Por outro, ela também pode possuir objetivos implícitos aprendidos, sem a interferência humana, diretamente a partir do conjunto de dados que recebe. Esses objetivos implícitos são difíceis de serem acompanhados, pois se baseiam em heurísticas e aproximações que podem gerar comportamentos inesperados, indesejados ou explicações inusitadas. Atualmente, muitas técnicas de aprendizagem não-supervisionadas têm se mostrado úteis para análise de grandes volumes de dados, ou modelagem de sistemas complexos (ex.: o jogo Go), tornando necessário desenvolver formas de se lidar com os objetivos implícitos.

Dentro deste contexto de inovação é importante saber se precisamos da IA. Se a resposta é afirmativa, é preciso ter claro qual IA queremos para o planeta (adotar uma perspectiva ambiental que não seja centrada apenas no humano), e quais as responsabilidades que ela terá. Questões como essas são fundamentais para se decidir como a regulamentação pode ocorrer para dar conta, pelo menos, no médio prazo, desta inovação.

As tecnologias da IA disponíveis atualmente podem ser classificadas em duas categorias de IA:

- que *aumenta o cérebro*, ou seja, aumenta as capacidades humanas, e
- que *imita o cérebro*, que o utiliza como uma inspiração biológica para seus modelos (termos utilizados na literatura internacional).

A IA que aumenta o cérebro é apresentada, pelas empresas de tecnologia, sempre com fins humanitários (a metáfora pode ser o uso de óculos, que aumenta a nossa habilidade de visão). Da mesma forma, vários dispositivos vestíveis da IA prometem aumentar alguma das capacidades humanas, como concentração e locomoção, proporcionando bem estar e qualidade de vida (ONU, 2023). Por exemplo, uma pessoa recebe um implante para recuperar sua capacidade de locomoção, porém a *startup* que forneceu o implante pode falir. O que acontecerá com a pessoa implantada? Retirar o implante, que não terá mais manutenção, e perder a tão sonhada capacidade de locomoção? Implantes podem melhorar a locomoção, mas uma vez que a tecnologia é conhecida, ela poderá também ser utilizada com outros fins. Se por um lado, essa tecnologia irá facilitar a vida de muitas pessoas, por outro, muitos produtos derivados

visam apenas aumentar a eficiência e a produtividade. Cada situação pode requerer uma regulação específica. Na falta de uma regulamentação geral os limites, propósitos e responsabilidades da tecnologia são desconhecidos. É difícil afirmar, a priori, se uma tecnologia será boa para todos.

Da mesma forma, os softwares de IA, em geral, prometem facilitar o dia a dia resumindo textos, gerando apresentações, vídeos, *podcast*, recomendando ou selecionando músicas de forma personalizada, indicando possíveis amigos e dizendo o que se gostaria de ter e comprar. No entanto, toda essa assistência tem um custo. Por exemplo, ao se fornecer um artigo recém escrito, para um aplicativo de IA que gera uma apresentação, para um evento científico, certamente, se ganha o tempo que seria necessário para produzir a apresentação e a IA estará ganhando o conhecimento contido no texto inédito. Este tipo de tecnologia já está sendo usada em várias áreas, como da programação (Gambacorta et al., 2024) e secretariado e administração (Weng, 2023). Além disso, o uso dos serviços de IA geralmente requer o pagamento de uma assinatura, que faz parte do modelo de negócio dessas tecnologias. Quando se recorre a estes recursos é importante conhecer, ao máximo, as suas regras. Alguns destes exemplos podem ser facilmente regulados. Basta que os softwares utilizados para os produzir adicionem uma marca que indique que os mesmos foram gerados pela IA. Neste cenário, o bom senso dirá se é uma situação onde o uso é aceitável ou não para aquele contexto específico. Como se observa, a IA não é exatamente uma amiga.

Mas, a IA não se resume à produtividade. Ela pode ser direcionada para a solução de alguns dos problemas da humanidade. O projeto Alpha Fold (Jumper *et al.*, 2021), por exemplo, se demonstrou capaz de prever a estrutura de proteínas, permitindo direcionar pesquisas com grande potencial para desenvolvimento novos medicamentos, assim como de alimentos mais nutritivos e que causem menos danos ambientais que a produção de proteína animal, beneficiando o planeta como um todo e não apenas aos humanos.

Do ponto de vista da pesquisa e do desenvolvimento da IA, destaca-se a IA que imita o cérebro para construir seus modelos. Aqui se enquadram as redes neurais, ou seja, tentativas de reproduzir, em algoritmos, a estrutura neural de organismos vivos inteligentes (servindo como inspiração histórica) e as arquiteturas neurais utilizadas para a programação de agentes proativos. Os agentes são uma categoria de sistemas de IA

que tomam decisões autônomas. Por exemplo, durante o dia, uma pessoa procura uma passagem para ir da cidade de Porto Alegre para São Paulo, mas considera o preço muito alto. O assistente pessoal de IA desta pessoa monitora suas tentativas de compra da passagem. À noite, os preços da passagem diminuem e o assistente pessoal aproveita a oferta e realiza o desejo de seu usuário. Essas tecnologias da IA se mostram eficientes na atualidade. Mas, também, trazem à tona questões importantes, como a necessidade de explicar as decisões tomadas por IAs com base nos seus objetivos (explícitos e implícitos) e a necessidade do desenvolvimento e do uso ético da IA.

ÉTICA E IA

A identificação de doutrinas e modelos éticos que orientam o comportamento humano tem sido um dos temas centrais do pensamento filosófico. Esta abordagem caracteriza a ética normativa, e é o ponto de relevância para o desenvolvimento da ética de máquina, ou seja, doutrinas humanas, mais uma vez, têm sido a base para a IA. Para Zoshak e Dew (2021), as doutrinas éticas mais aplicadas no desenvolvimento de sistemas de IA são a Deontologia (cumprimento de deveres e obrigações morais) e o Utilitarismo (geração do maior bem ou felicidade para o maior número de pessoas).

A ética deontológica sustenta que uma dada ação deve ser julgada com base na sua compatibilidade com um conjunto de deveres reconhecidos como legítimos pelos tomadores de decisões racionais (Vamplew *et al.*, 2018). Estruturas baseadas nesta doutrina são orientadas a deveres (Pfordten *et al.*, 2012). Esta estrutura vincula o julgamento moral à intenção de quem age. Por essa razão, também pode ser chamada de ética intencionalista. Os modelos proativos das IAs que temos atualmente são baseados em objetivos a serem alcançados, sejam eles implícitos ou explícitos, ou seja são intencionais.

Na abordagem utilitarista, assume-se que a desejabilidade de uma ação pode ser medida por métricas de utilidade e que uma ação é julgada moralmente correta se as suas consequências levam à melhor utilidade possível (Vamplew *et al.*, 2018). No entanto, existem diferentes perspectivas quanto à definição de utilidade que se pretende maximizar, ou seja, não há consenso sobre qual a maneira correta de combinar escalas de utilidade ou mesmo se é conveniente combiná-las. O problema da utilidade é que o seu cálculo pode conter heurísticas indesejáveis (ex.: racismo algorítmico). Portanto, a utilidade talvez não seja uma teoria adequada para abordar as questões éticas da IA.

ESTADO DA ARTE

Desde 2009, o campo da IA tem sido objeto de intensas discussões sobre sua regulamentação (Asaro, 2012), com campanhas visam banir o uso de Armas Autônomas Letais (Lethal Autonomous Weapons, ou LAWs), sistemas robóticos que tenham como função ferir ou eliminar seres humanos autonomamente. Essa posição é mantida desde 2018 pelas Nações Unidas (United Nations, 2023). Na década de 2020, com os avanços em *deep learning* e o lançamento da IA Gerativa para o grande público, há uma ampliação e uma mudança no direcionamento dessas discussões. Governos e organizações não governamentais passaram a publicar *guidelines* para a pesquisa e o desenvolvimento da IA, além de propor legislações voltadas à sua regulamentação. Casos como a greve do sindicato dos escritores de Hollywood em 2023 evidenciou que o impacto dessa tecnologia já chegou à sociedade, e questões relacionadas ao trabalho também estão sendo afetadas, destacando a urgência de se discutir suas implicações em diferentes esferas da sociedade.

Nos últimos anos, aproximadamente oitenta nações publicaram guias e estratégias nacionais voltadas ao desenvolvimento da IA (AI Index, 2024), assim como vários documentos emitidos pela ONU - classificados como “soft law” - instrumentos normativos de direito internacional que não possuem força legal, sanções ou caráter vinculativo, mas com potencial de regular comportamentos sociais. A exemplo da Estratégia Brasileira de Inteligência Artificial (Brasil, 2021), que destaca as principais questões éticas e define prioridades de investimento relacionadas a essa tecnologia. Mas essas iniciativas representam apenas um primeiro passo em direção à regulação e precisam ser coordenadas entre si. Alcançar tal coordenação, no entanto, é um desafio significativo.

Até o momento, a única legislação em vigor que regulamenta a IA é o “AI Act” da União Europeia, cujo objetivo é garantir que os sistemas de IA usados na UE sejam seguros, transparentes, rastreáveis, não-discriminatórios e ecológicos. Sistemas de IA devem ser supervisionados por pessoas, em vez de operarem de forma completamente automatizada, para prevenir resultados prejudiciais” (European Union, 2023). A legislação europeia não estabelece direitos de usuários, ou tipifica crimes que as plataformas (aplicações da IA) podem causar, mas classifica os sistemas de acordo com seus níveis de risco, estabelecendo diretrizes de como eles devem funcionar.

A maior parte de seu desenho é baseado em casos, pois determina o seu grau de periculosidade a partir do uso do sistema. Por exemplo, o uso sistemas que utilizam técnicas subliminares que vão além da capacidade da consciência de um indivíduo, ou que aplicam técnicas propositadamente manipulativas ou enganosas, são considerados de risco extremo e, portanto, proibido. Cabe destacar que países como a China, a Índia e o Brasil Lei 8/2023 (Brasil, 2023), dentre outros, estão em diferentes fases de desenvolvimentos de suas legislações, mas todos, em certa medida, inspirados na legislação europeia.

Ainda, é complexo estabelecer uma dinâmica entre os processos regulatórios locais, regionais e globais. Esse fator pode contribuir para que o avanço da IA reproduza ou amplie as desigualdades tecnológicas e econômicas já existentes. Outro aspecto de difícil gestão é o fato de que as legislações não possuem natureza flexíveis e esse aspecto entra em choque com o avanço da pesquisa em áreas dinâmicas, como a tecnologia. Dessa forma, uma regulação mais flexível e adaptativa desde a origem pode ser uma abordagem mais adequada.

MUDAR O PONTO DE PARTIDA

O aprendizado de máquina, especialmente quando se baseia em grandes volumes de dados (*big data*), trouxe à tona questões éticas para a coleta, curadoria e armazenamento dos dados, além de levantar preocupações sobre os algoritmos que dificultam a explicação (objetivos implícitos e pesos autorregulados) de suas previsões ou predições, recomendações e criações. Sistemas de *deep learning*, como comentado anteriormente, usam redes neurais complexas/profundas, com muitas camadas intermediárias, mostram que o potencial de um modelo de IA pode estar inversamente relacionado à sua *explicabilidade*. Ou seja, a ética precisa estar incorporada tanto no uso quanto no seu desenvolvimento da IA, abrangendo dados, algoritmos e modelos.

Incluir ética no projeto, desenvolvimento e teste de sistemas de IA não é uma tarefa simples. A dificuldade reside em escolher referenciais teóricos e modelos que possam ser transpostos para as várias etapas do ciclo de vida dos sistemas de IA e trabalhar com equipes multidisciplinares durante todo o processo, o que é conhecido por “ética por design”. Vale destacar que, em IA, fala-se em previsões e não em resultados. Isso porque a IA é capaz de trabalhar com falta de informações e modelos não determinísticos, o que introduz um segundo desafio ético: a responsabilidade. Ou seja, as previsões podem ser corretas, parcialmente corretas ou incorretas. Ainda, várias

possibilidades de previsões ou criações podem ser fornecidas a partir de uma única consulta, com diferentes probabilidades de estarem ou não corretas. Mesmo com estas ressalvas, existem boas práticas prévias, como o uso do padrão ISO (Organização Internacional de Normalização) para o desenvolvimento de software, que podem servir de inspiração para a IA.

É necessário também educar as pessoas para fazerem o uso ético e consciente da IA. Uma cultura sólida de uso poderia atenuar a necessidade de regulação, permitindo aos usuários se protegerem de potenciais perigos e abusos. Essa abordagem faz uma diferença significativa, pois muda o ponto de vista. Ao invés de buscar regular casos ou aplicações específicas, que como visto nos vários exemplos, apresentam graus variados de complexidade e necessidades distintas. Ainda, regular aplicações da IA é uma tarefa desafiadora, pois seu desenvolvimento é contínuo e as aplicações alcançam uma vasta gama de atividades humanas. Também é importante se ter em mente que a regulação pode engessar a pesquisa na área.

Portanto, regular a IA é difícil sob o ponto de vista das suas aplicações, pois a legislação pode vir tarde demais. No entanto, existe uma alternativa viável: repensar a regulação da IA a partir de seus elementos fundamentais, ou seja, os dados, algoritmos e modelos. Isso inclui:

- onde os dados são coletados, como são coletados, para que estes dados são coletados, como é realizada a curadoria e onde serão armazenados;
- quem constrói os algoritmos, com que dados são treinados, possuem objetivos implícitos, quais são os objetivos explícitos, possui aprendizado em tempo real, como são tomadas as decisões, se é possível identificar o conjunto de dados utilizado, pelo algoritmo, para uma decisão, em particular;
- qual a finalidade de um modelo. Ele é benéfico para os humanos? quais os seus potenciais impactos sociais, econômicos e ambientais?

Uma legislação que parte desses pilares da IA e que trate possíveis riscos na origem (seu desenvolvimento através de boas práticas) pode ser mais abrangente e sustentável do que aquela que foca apenas na regulamentação dos efeitos da IA.

Um exemplo relevante dessa abordagem é o Solid, um projeto idealizado por Tim Berners-Lee (o mesmo que inventou a *World Wide Web*), um projeto de código aberto que dá mais controle sobre o armazenamento e tratamento dos seus dados na web. Berners-Lee começou a reconhecer que talvez a Internet atual não seja a sua

melhor versão e propôs uma forma nova de descentralizar os dados, permitindo que cada indivíduo receba um armazenamento de dados pessoais *on-line*. Ou seja, com o Solid, os usuários podem decidir quais dados serão armazenados e quem terá acesso a essa área pessoal, garantindo uma maior proteção e controle sobre a privacidade desde a origem dos dados. Essa iniciativa contempla, em parte, a privacidade dos dados na origem. Um exemplo de aplicação concreta do Solid ocorre na Bélgica que o utiliza para compartilhar informações e comunicações médicas, garantindo a segurança dos dados pessoais dos cidadãos. Esse modelo proposto inicialmente por Berners-Lee, descentraliza o controle sobre os dados pessoais e reforça a importância de uma regulamentação a partir do desenvolvimento tecnológico e não apenas em suas consequências.

No artigo 50 da legislação da União Europeia é possível ter alguns vislumbres de uma legislação voltada aos pilares da IA, que traz obrigações de transparência, por exemplo, informando o usuário quando ele estiver lidando com um sistema artificial. Nessa abordagem, a transparência é uma regulação a ser feita sobre o modelo da IA, garantindo que o usuário saiba quando está, ou não, lidando com uma IA. Quando o enfoque recai sobre os pilares da IA, a legislação não se baseia em casos e aplicações, mas em mitigação de riscos inerentes à natureza da IA. Para regular, é preciso primeiro entender o que se está regulando. E, de modo geral, há uma compreensão limitada sobre a IA.

Para ilustrar esta proposta, utilizaremos o exemplo da ética. A ética é necessária, tanto no desenvolvimento, quanto no uso da IA, permeando a coleta e o uso dos dados e no desenvolvimento dos algoritmos que utilizam objetivos explícitos e implícitos. No primeiro caso, utilizar princípios éticos depende dos desenvolvedores humanos. No segundo, depende do processo de coleta, curadoria e armazenamento dos dados. Nessas duas etapas a educação para o desenvolvimento ético da IA é fundamental.

Outra necessidade é estabelecer regras para as empresas que desenvolvem estes produtos. Por exemplo, poderia ser adotado um sistema de certificação semelhante ao selo "produto orgânico", indicando se um modelo de IA atende a critérios específicos e recebendo um selo de "IA Saudável". No caso dos modelos, como eles são o resultado dos dados e dos algoritmos, eles já estariam adequados ao recebimento do selo de "*IA saudável*". A outra parte da equação vai depender do uso ético desses modelos, pelas pessoas. Os produtores da IA poderiam exibir o selo e caberia a nós optarmos por uma IA saudável ou não.

Com medidas de regulação centradas na IA o Brasil poderia, por exemplo, priorizar a compra dos produtos conhecidos como AI4SG (*Artificial Intelligence for Social Good*, ou em tradução direta, Inteligência Artificial para o Bem Social) (Floridi, L. *et al.*, 2020), e coibir o uso de sistema de repressão, discriminatórios e preconceituosos que aumentariam as disparidades sociais, econômicas e ambientais. No contexto brasileiro, é difícil controlar este processo, uma vez que as *big techs* não estão sediadas no país, porém o Brasil pode utilizar princípios éticos ao que é desenvolvido domesticamente e pressionar os organismos internacionais para que tratem os efeitos colaterais da IA de forma unificada. A pressão pode ocorrer de forma solitária ou acompanhada em conjunto com organismos ou blocos, como por exemplo, o Mercosul, G20 e Brics. O Brasil pode também proteger dados estabelecendo que, independentemente do produto de IA que se esteja utilizando, os dados dos brasileiros permaneçam no país. Essa medida é particularmente importante para aplicações da IA que envolvam a segurança, saúde e a educação. Existem várias maneiras de se emitir um selo para a IA. Defendemos a adoção de um sistema semelhante ao utilizado para os produtos orgânicos, que é a verificação por pares.

Atualmente em análise no congresso brasileiro, a Lei 2338/2023 (Brasil, 2023), que visa regular as IAs no país, discriminando os agentes (operadores e fornecedores de sistemas de IA) e suas responsabilidades. O projeto, em acordo com a Lei Geral de Proteção de Dados (Brasil 2018) tem uma forte preocupação em garantir as proteções sobre os dados pessoais, de forma similar ao estabelecido na legislação europeia. A proposta legislativa classifica tipos de usos inaceitáveis pelos riscos, e exige a adoção de sistemas de governança da IA, especialmente para mitigação de riscos. Apesar dessa sistemática estar descrita, observa-se que o projeto não define a palavra algoritmo e modelo, apenas exigindo que se tomem medidas para garantir a *explicabilidade* dos sistemas. Isso acaba tornando a legislação vaga no quesito de desenvolvimento e à aplicação, uma vez que a proposta sequer regula como as bases de dados de treinamento podem ser obtidas.

O Brasil perdeu a corrida IA desde os anos 90. Segundo Chakravorti *et al.* (2021), a concentração de especialistas em IA no Brasil é de 0,2, em uma escala de [0-10]. Até o momento, o país não teve uma política abrangente de formação de talentos para áreas consideradas estratégicas pelo CNPq (Conselho Nacional de Desenvolvimento Científico e Tecnológico).

Logo, o Brasil não é um ator mundial no que se refere a esta tecnologia (Adegoke *et al.*, 2024). Isso gera uma situação delicada, pois nos torna, predominantemente, usuários e também, possivelmente vulnerável ao desemprego ou subemprego em escala global, visto que, cada vez mais é necessária formação para se compreender e desenvolver tecnologias inovadoras. Isso gera o risco da dominação geopolítica mediada pela tecnologia.

EDUCAÇÃO PARA A IA: O COMPORTAMENTO ÉTICO VOLUNTÁRIO

A proposta de repensar a regulação da IA fundamenta-se em dois aspectos: através do selo de qualidade de origem, mas também, da necessidade de se educar as pessoas para o consumo saudável da IA. Ou seja, como as pessoas utilizam os modelos de IA. A IA é uma inovação e é uma realidade contemporânea que facilita diversos setores das nossas vidas, pois pode aumentar nossas habilidades. A tendência é que ela seja cada vez mais difundida. Isso inclui alunos, professores e administradores escolares. Ademais, como visto, é difícil controlar a IA e o acesso à IA (ex.: smartphones já possuem uma IA embarcada por padrão). Logo, educar para o uso ético e consciente da IA pode ser uma das melhores formas de se lidar com ela. A IA afeta a forma como se ensina, afeta os empregos e a renda dessas gerações, afeta as relações sociais (como as pessoas se comunicam) e afeta o meio ambiente.

Essa realidade exige muito cuidado e preparo, mostrando a importância da escola para transformar o que ainda é possível nessa conjuntura. A inovação requer novas habilidades e competências. Caso não sejam desenvolvidas, primeiramente os empregos serão perdidos para quem souber utilizar a inovação para aumentar suas capacidades. As instituições de ensino podem ajudar neste contexto. Enquanto grupo de pesquisa em Inteligência Artificial aplicada à Educação, defendemos que é necessário abordar a IA de duas formas: Pensar **com** a IA, que consiste em usar a IA para auxiliar na resolução de problemas, e Pensar **sobre** a IA, que implica em entender como a IA resolve os problemas (Referência omitida para submissão anônima). Essas duas dimensões contribuem para um uso voluntário que seja consciente e ético, pois não é possível compreender os impactos sem compreender o funcionamento da tecnologia.

É importante, também, que a escola desenvolva habilidades e competências em seus alunos, para que eles estejam preparados para enfrentar as mudanças advindas da inovação. Profissões tradicionais estão desaparecendo ou sendo profundamente transformadas; habilidades e competências até então fundamentais o deixam de ser

enquanto outras se tornam necessárias. Por exemplo, até pouco tempo, a firmeza nos movimentos era uma habilidade fundamental para um cirurgião. Atualmente as cirurgias são realizadas com o apoio de braços robóticos, que são mais precisos que os humanos. Ter adquirido a habilidade de movimentar *joystick* pode ser muito útil nesse contexto. Usando esse mesmo exemplo, é necessário cuidado com a governança da IA. Supondo que o sindicato médico, de um país, mova uma ação visando a proibição da análise de exames de imagens ser realizada por IA, determinando que apenas médicos podem realizar essa tarefa. Sabe-se que a IA é muito eficiente nesta área, alcançando até 87% de acerto. O objetivo do sindicato é o de proteger esses empregos (ação corporativa). Como agiria o regulador? levaria em consideração o bem estar comum, o fato do custo da análise realizada pela IA ser menor para as pessoas ou, preferirá preservar os empregos desse segmento? Essa situação pode ocorrer em muitas outras áreas da atividade humana atual.

Estabelecer o conjunto mínimo de habilidades e competências que a escola, de hoje, precisa desenvolver em seus alunos para eles lidarem com a possível perda do emprego para a tecnologia é uma tarefa difícil, mas que precisa ser abordada.

Ainda, a genIA (*Generative AI*, ou em tradução livre IA Gerativa), que não é nada ecológica devido ao seu alto consumo de energia e água (podendo não ser adequada para todos os seres vivos), intensificou essas implicações, em particular na educação. Se por um lado, quando bem utilizada, pode aumentar as habilidades humanas, por outro, está levando áreas da atividade humana a se repensarem ou mesmo deixando de ser economicamente viáveis (ex.: designers, ilustradores, tradutores, entre outros). A genIA usa as duas linhas de inovação da IA: a IA que aumenta as capacidades do cérebro e a IA que imita o cérebro, pois utiliza redes neurais para o seu treinamento. Ela gera conteúdos baseados no que aprende de produções prévias geradas por humanos. Mas, isso é a genIA na sua infância e, mesmo assim, já traz muitas questões para a educação - por exemplo, de quem é a autoria de uma produção se a IA for usada apenas para traduzir um texto de autor? Pode-se ter uma máquina como coautora de um trabalho acadêmico? Estas e outras questões, relacionadas com a integridade acadêmica, que a educação precisa abordar no momento. Ainda, a genIA gera suas produções apenas com base em probabilidades, como por exemplo, a probabilidade da palavra 'x' ser seguida pela palavra 'y'. Não existe propósito intrínseco nas produções e, mesmo assim, já causa um considerável debate. Pode-se

imaginar o que está por vir: em breve a IA poderá ser um agente dotado de raciocínio, conhecimento e poder de decisão, ou seja, autonomia.

Retomando a questão da educação, neste caso continuada ao longo da vida, a sociedade atual preocupa-se com requalificação para o trabalho necessária tanto pelo avanço da inovação, quanto pelo aumento da expectativa de vida das pessoas, e a consequente necessidade de se continuar ativos e com renda. É necessário assegurar que o modelo tecno-econômico não se torne autofágico. A sociedade pode tratar dessas situações considerando as perguntas iniciais deste artigo: "Precisamos da IA?", "É possível fazer diferente?", regulando a IA adequadamente para que seja possível aproveitar os seus benefícios e evitar os seus efeitos colaterais. Governos democráticos podem agir para que o controle da IA seja guiado pelo bem comum do planeta e não apenas pelos interesses de apenas algumas empresas que produzem essa tecnologia.

A pesquisa também pode ser direcionada para avançar nos algoritmos de aprendizado, com base em pequenas quantidades de dados, buscando a otimização do treinamento, que atualmente gasta uma grande quantidade de tempo de processamento, energia e água. Reduzir e refinar as bases de treinamento, assim como mitigar os efeitos colaterais dos modelos, são alguns aspectos nos quais a pesquisa pode contribuir. Portanto, a regulação precisa encontrar a medida certa para não inibir o desenvolvimento científico.

CONSIDERAÇÕES FINAIS

A Inteligência Artificial é uma tecnologia disruptiva que perturba os modelos econômicos, sociais e ambientais vigentes e induz a formação de outros e novos paradigmas. Cabe a nós avaliar as opções disponíveis e decidir qual modelo se deseja seguir, considerando se o modelo escolhido atende às necessidades do planeta. Certamente, o modelo atual tem a tendência de contemplar a concentração de renda e poder para alguns países que dominam esta tecnologia, visto que a IA representa conhecimento, e conhecimento é poder.

Mas, do que irão viver as pessoas do restante do mundo? De salário social? Os habitantes dos países que não geram a riqueza tecnológica terão uma mesma renda mínima universal sustentada por impostos cobrados de produtos tecnológicos, (neste caso específico de robôs e da IA) e isso constituirá a equidade social? Mas haverá um bolo? Como será a divisão desse bolo? Nem o proposto por Domenico De Masi (2000), no seu livro “O ócio criativo” parece não ser mais possível. Até a chegada da IA

Gerativa, era possível acreditar que os humanos, em um mundo ideal, poderiam se dedicar às atividades criativas, enquanto as demais poderiam ser ocupadas pela tecnologia. Essa ideia parecia paradisíaca. Ainda pode ser?

E o que dizer do fato de que as produções da IA podem gerar o efeito da “rolagem infinita”, ou seja, onde os dados gerados pela IA, nem sempre verdadeiros, sejam utilizados para realimentar a própria IA, alcançando assim sua autonomia na produção dos dados necessários para o seu contínuo treinamento. Diante desses desafios, torna-se necessário aprofundar o debate sobre a regulação ou a não regulação da IA, bem como investir em educação e pesquisa que promovam uma compreensão mais ampla de suas implicações. Como se pode perceber, este texto não apresenta conclusões, e sim considerações ainda com perguntas que ainda necessitam de respostas.

REFERÊNCIAS BIBLIOGRÁFICAS

ADEGOKE, Adeboye *et al.* **Governança da inteligência artificial (IA): a interação entre o local, o regional e o global em direção a uma solidariedade transnacional** | **Revista poliTICs**. [S. l.], 2024. Disponível em: <https://politics.org.br/pt-br/outros-assuntos-news/governanca-da-inteligencia-artificial-ia-interacao-entre-o-local-o-regional-e>. Acesso em: 19 out. 2024.

AI INDEX. **Countries with national artificial intelligence strategies**. [S. l.], 2024. Disponível em: <https://ourworldindata.org/grapher/national-strategies-on-artificial-intelligence>. Acesso em: 18 out. 2024.

ASARO, Peter. On banning autonomous weapon systems: human rights, automation, and the dehumanization of lethal decision-making. **International Review of the Red Cross**, [s. l.], v. 94, n. 886, p. 687–709, 2012. Disponível em: https://www.cambridge.org/core/product/identifier/S1816383112000768/type/journal_article. Acesso em: 18 out. 2024.

BRASIL. **Estratégia Brasileira de Inteligência Artificial (EBIA)**. Brasília, DF: Ministério da Ciência, Tecnologia e Inovações, 2021. Disponível em: https://www.gov.br/mcti/pt-br/acompanhe-o-mcti/transformacaodigital/arquivosinteligenciaartificial/ebia-documento_referencia_4-979_2021.pdf. Acesso em: 17 out. 2024.

BRASIL. **Lei 13709/2018**. 2018. Disponível em: https://www.planalto.gov.br/ccivil_03/_ato2015-2018/2018/lei/113709.htm. Acesso em: 19 out. 2024.

BRASIL. **PL 2338/2023**. [S. l.], 2023. Disponível em: <https://www25.senado.leg.br/web/atividade/materias/-/materia/157233>. Acesso em: 19 out. 2024.

CHAKRAVORTI, Bhaskar *et al.* 50 Global Hubs for Top AI Talent. **Harvard Business Review**, [s. l.], 2021. Disponível em: <https://hbr.org/2021/12/50-global-hubs-for-top-ai-talent>. Acesso em: 19 out. 2024.

DAVID, Mylany; ROBITAILLE, Miriam; GAGNON, Samuel. Automatisation dans la rédaction des actes juridiques : usurpe-t-on les fonctions des juristes? *In*: LANGLOIS AVOCATS. 4 ago. 2017. Disponível em: <http://langlois.ca/ressources/automatisation-dans-la-redaction-des-actes-juridiques-usurpe-t-les-fonctions-des-juristes/>. Acesso em: 15 out. 2024.

EUROPEAN UNION. **EU AI Act: first regulation on artificial intelligence**. [S. l.], 2023. Disponível em: <https://www.europarl.europa.eu/topics/en/article/20230601STO93804/eu-ai-act-first-regulation-on-artificial-intelligence>. Acesso em: 19 out. 2024.

EUROPEAN UNION. **EU Artificial Intelligence Act**. [S. l.: s. n.], 2021. Disponível em: <https://artificialintelligenceact.eu/>. Acesso em: 19 out. 2024.

FLORIDI, Luciano *et al.* How to Design AI for Social Good: Seven Essential Factors. **Science and Engineering Ethics**, [s. l.], v. 26, n. 3, p. 1771–1796, 2020. Disponível em: <https://doi.org/10.1007/s11948-020-00213-5>. Acesso em: 19 out. 2024.

GAMBACORTA, Leonardo *et al.* Generative AI and labour productivity: a field experiment on coding. [s. l.], 2024. Disponível em: <https://www.bis.org/publ/work1208.htm>. Acesso em: 15 out. 2024.

JUMPER, John *et al.* Highly accurate protein structure prediction with AlphaFold. **Nature**, [s. l.], v. 596, n. 7873, p. 583–589, 2021. Disponível em: <https://www.nature.com/articles/s41586-021-03819-2>. Acesso em: 15 out. 2024.

MASI, Domenico De. **O Ocio Criativo**. [S. l.]: Sextante, 2000.

ONU. **Governing AI for Humanity: Interim Report / United Nations AI Advisory Body**. [S. l.]: United Nations, 2023. Disponível em: [https://www.un.org/sites/un2.un.org/files/un_ai_advisory_body_governing_ai_for_hum anity_interim_report.pdf](https://www.un.org/sites/un2.un.org/files/un_ai_advisory_body_governing_ai_for_hum_anity_interim_report.pdf).

SALOMÃO, Luis Felipe. Artificial intelligence: technology applied to conflict management within the Brazilian judiciary. [s. l.], 2022. Disponível em: <https://hdl.handle.net/10438/33954>. Acesso em: 15 out. 2024.

SCHNEIER, Bruce; SANDERS, Nathan. **Six ways that AI could change politics**. [S. l.], 2023. Disponível em: <https://www.technologyreview.com/2023/07/28/1076756/six-ways-that-ai-could-change-politics/>. Acesso em: 15 out. 2024.

UNITED NATIONS. Lethal Autonomous Weapon Systems (LAWS) – UNODA. *In*: [s. d.]. Disponível em: <https://disarmament.unoda.org/the-convention-on-certain-conventional-weapons/background-on-laws-in-the-ccw/>. Acesso em: 18 out. 2024.

VAMPLEW, Peter *et al.* Human-aligned artificial intelligence is a multiobjective problem. **Ethics and Information Technology**, [s. l.], v. 20, n. 1, p. 27–40, 2018. Disponível em: <http://link.springer.com/10.1007/s10676-017-9440-6>. Acesso em: 15 out. 2024.

Referência omitida para submissão anônima.

VON DER PFORDTEN, Dietmar. Five Elements of Normative Ethics - A General Theory of Normative Individualism. **Ethical Theory and Moral Practice**, [s. l.], v. 15, n. 4, p. 449–471, 2012. Disponível em: <https://link.springer.com/10.1007/s10677-011-9299-2>. Acesso em: 15 out. 2024.

WENG, Jiaxiong (Connor). **Putting Intellectual Robots to Work: Implementing Generative AI Tools in Project Management**. [S. l.]: NYU SPS Applied Analytics Laboratory, 2023. Technical Report. Disponível em: <http://archive.nyu.edu/handle/2451/69531>. Acesso em: 15 out. 2024.

ZOSHAK, John; DEW, Kristin. Beyond Kant and Bentham: How Ethical Theories are being used in Artificial Moral Agents. *In*: CHI '21: CHI CONFERENCE ON HUMAN FACTORS IN COMPUTING SYSTEMS, 2021, Yokohama Japan. **Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems**. Yokohama Japan: ACM, 2021. p. 1–15. Disponível em: <https://dl.acm.org/doi/10.1145/3411764.3445102>. Acesso em: 15 out. 2024.