



SENADO FEDERAL
SECRETARIA-GERAL DA MESA
SECRETARIA DE REGISTRO E REDAÇÃO PARLAMENTAR
REUNIÃO

19/10/2023 - 7ª - Comissão Temporária Interna sobre Inteligência Artificial no Brasil

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO. Fala da Presidência.) - Declaro aberta a 7ª Reunião da Comissão Temporária Interna sobre Inteligência Artificial no Brasil, criada pelo Requerimento nº 722, de 2023, com a finalidade de, no prazo de até 120 dias, examinar os projetos concernentes ao relatório final aprovado pela Comissão de Juristas, responsável por subsidiar a elaboração do substitutivo sobre inteligência artificial no Brasil, criada pelo Ato do Presidente do Senado Federal nº 4, de 2022, bem como eventuais novos projetos que disciplinem a matéria.

Pauta.

Informo que esta reunião se destina à realização de audiência pública com a finalidade de debater aspectos gerais da inteligência artificial, com vistas a compreender, quanto ao espectro jurídico, justificativas, definições, afetação a direitos fundamentais, princípios e fundamentos da regulação; e, no tocante ao espectro técnico, os tipos, modelagens e diferentes aplicações da tecnologia, em cumprimento ao Requerimento nº 4, de 2023-CTIA, de minha autoria.

O público interessado em participar desta audiência pública poderá enviar perguntas ou comentários pelo endereço www.senado.leg.br/ecidadania ou ligar para 0800 0612211.

Encontram-se presentes no plenário da Comissão, além do nosso querido Vice-Presidente, Senador Marcos Pontes: Lucas Borges de Carvalho, Gerente de Projetos da Autoridade Nacional de Proteção de Dados (ANPD); Gustavo Zaniboni, Presidente da Coordenação de Inteligência Artificial da Associação Brasileira de Governança Pública de Dados Pessoais (govDADOS). Encontram-se também presentes, por meio de sistema de videoconferência: Diogo Cortiz, Coordenador do Mestrado e Doutorado em Tecnologias da Inteligência e Design Digital da Pontifícia Universidade Católica de São Paulo (PUC-SP); Marcelo Finger, Professor Titular do Departamento de Ciência da Computação do Instituto de Matemática e Estatística da Universidade de São Paulo (USP); André Carlos Ponce de Leon Ferreira de Carvalho, Diretor do Instituto de Ciências Matemáticas e de Computação da Universidade de São Paulo (ICMC-USP), representando Thais Vasconcelos Batista, Presidente da Sociedade Brasileira de Computação (SBC); Fernando Malerbi, Coordenador do Departamento de Saúde Ocular da Sociedade Brasileira de Diabetes (SBD); Nina da Hora, Diretora Executiva do Instituto da Hora; André Lucas Fernandes, Diretor e Fundador do Instituto de Pesquisa em Direito e Tecnologia do Recife.

Agradeço a presença de todos.

(Intervenção fora do microfone.)

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO) - Eu gostaria de pedir aos dois que estarão presentes aqui na mesa, Lucas Borges de Carvalho e Gustavo Zaniboni, se quiserem fazer parte já da mesa, são nossos convidados.

E gostaria de ver se já está em condições de fazer sua participação virtual Fernando Malerbi, que nos pediu essa participação um pouco antes, por conta de um compromisso. Então, nós vamos compreender aqui a adequação de horário.

Eu passo, então, a palavra ao Dr. Fernando Malerbi, Coordenador do Departamento de Saúde Ocular da Sociedade Brasileira de Diabetes. V.Sa. tem a palavra por até dez minutos para sua intervenção.

O SR. FERNANDO MALERBI (Para expor. *Por videoconferência.*) - Prezado Sr. Senador, muito obrigado pela palavra. A minha câmera está desabilitada aqui pelo anfitrião, mas eu estou vendo todos e espero que estejam me ouvindo. *(Pausa.)*

Pronto, agora sim.

Muito obrigado pela oportunidade de representar a Sociedade Brasileira de Diabetes nesta importante discussão.

Eu vou ser breve aqui na minha fala.

A gente, como sociedade de diabetes, tem como missão contribuir sempre para a prevenção e o tratamento adequado do diabetes e suas complicações. É uma doença que acomete milhões de indivíduos no Brasil. E a gente tem como missão também disseminar conhecimento técnico, científico entre médicos e profissionais de saúde e a conscientização da população sobre essa doença, para melhorar a qualidade de vida das pessoas com diabetes, sobretudo, colaborando com o Estado na formulação e execução de políticas públicas voltadas para a atenção correta dos pacientes e para a redução significativa da carga dessa doença que afeta tantos brasileiros.

Nesse âmbito, nós vemos com muito bons olhos a incorporação de tecnologia que possa realmente ampliar o acesso dos pacientes aos cuidados com o diabetes e à prevenção de suas complicações. Nesse âmbito, nós temos um grande interesse e atuação em colaborar nessa pauta de incorporar tecnologia de inteligência artificial no cuidado da pessoa com diabetes.

Eu queria brevemente mencionar quais são as possibilidades atuais do emprego da IA em diabetes. Em primeiro lugar, para aqueles indivíduos que usam insulina para o tratamento de sua doença, a IA tem mostrado um fator muito importante em estudos para controlar a dose do medicamento em sistemas artificiais e também para dar suporte à decisão, para o indivíduo e para a equipe de saúde, em termos do melhor tipo de tratamento.

Adicionalmente, a IA tem se mostrado muito importante no rastreamento de complicações crônicas da diabetes, sobretudo a complicação crônica ocular. A retinopatia diabética é uma das mais importantes causas de cegueira evitável na população adulta. Então, é de suma importância que a gente consiga rastrear os indivíduos com diabetes, fazer uma detecção dessa doença nas fases iniciais e propiciar um tratamento adequado, com os oftalmologistas, para evitar que essa doença se agrave e leve à perda da visão, que é o que infelizmente ocorre quando uma pessoa com diabetes não tem o acesso, em tempo adequado, a um especialista, que pode, de fato, fornecer o tratamento.

Além da complicação ocular, temos também a possibilidade de emprego da IA na detecção de feridas no pé do indivíduo com diabetes, o que consiste infelizmente numa grande causa de amputação em nosso meio. Detectando então inicialmente feridas e dando um tratamento adequado, nós conseguimos melhorar muito esse desfecho e evitar amputações de membros inferiores.

E por último, o próprio paciente que lida com a doença, que tem uma rotina já dificultada por ter que ver as atividades diárias, a questão de dieta e tratamento e visitas aos médicos, pode se beneficiar, por exemplo, com aplicativos que, através de uma fotografia de um prato de alimentos, vai revelar o conteúdo nutricional e vai permitir que ele regule a dose de tratamento necessário.

Então, com essa introdução, eu queria dizer que a IA depende de uma qualidade robusta de uma base de dados. A gente entende que os produtos devem ter esse alicerce, esse fundamento numa base de dados que tenha representatividade, pensando em evitar risco de viés, e ter uma amostra que permita, a partir dela, a construção de ferramentas.

Em termos de implementação no mundo real, na linha do que o relatório final da Comissão de Juristas mostrou, nós acreditamos que, na área da saúde em geral, vai haver uma regulação setorial. A gente esteve em contato, inclusive, com técnicos da Anvisa, no ano passado, em Brasília, para discutir esse assunto. Fomos muito bem recebidos por técnicos de excelente gabarito. Mas não está claro, em termos legais, como vai se dar essa regulação em vários aspectos. Por exemplo, se vai haver diferença na regulação de algoritmos autônomos, ou seja, naqueles que dão uma informação automaticamente ou naqueles assistivos, ou seja, aqueles que podem auxiliar um médico ou um profissional de saúde na decisão clínica. Então, existe essa diferença.

Também há a importância da classe de risco. Vejam os senhores que o risco de um algoritmo errar na hora em que ele prevê o componente nutricional de um prato, de um almoço, é diferente de ele errar na predição de que se tem ou não tem uma doença na retina, no fundo do olho, que possa levar à cegueira e, sobretudo, no tipo de tratamento de insulina que vai ser empregado. Então, cada aplicação tem o seu risco respectivo e a regulação, certamente, tem que contemplar essas classes de risco.

É importante entender, também, como que, para o usuário final, essa informação gerada pelo sistema automático vai ser apresentada, se é um dado numérico, se é uma resposta binária tipo "sim" ou "não", se é um mapa de calor com imagens realçadas. Isso também seria motivo de debate, de estudo.

No próprio caminho regulatório, de acordo com a classe de risco, como desenvolvedores podem ter um caminho facilitado para isso não coibir a pesquisa e o avanço que podem beneficiar os pacientes, mas também com um olhar responsável das autoridades de liberar com critérios muito bem definidos.

Há necessidade de validação externa de algoritmos que estejam sendo desenvolvidos, ou seja, existe uma prova de que ele, de fato, funciona numa população diferente daquela na qual ele foi concebido. E, sobretudo, há a questão de educação dos usuários, de todos os envolvidos, tanto dos profissionais de saúde que vão adotar esse recurso para uma melhor qualidade de vida dos seus pacientes, quanto do próprio usuário final para, sobretudo, desfazer crenças infundadas que podem evitar a implementação desse tipo de tecnologia no mundo real. E a gente está falando, então, de atividades educativas que podem ser capitaneadas tanto pelo Estado quanto por desenvolvedores e sociedades médicas. A academia tem um papel muito importante. Há a questão da responsabilização, em termos de situações em que o algoritmo entrega um resultado que não condiz com o que seria melhor, prático. Enfim, são diversos temas.

Eu queria apenas introduzir esses temas todos para o debate e dizer que a Sociedade Brasileira de Diabetes está à disposição da sociedade brasileira em geral para auxiliar nessas pautas. Nós temos um corpo bastante gabaritado, com profissionais especializados em várias áreas que dizem respeito à atenção, ao cuidado com o indivíduo com diabetes. Estamos aqui à disposição para contribuir neste debate.

Muito obrigado.

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO) - Queremos agradecer essa participação importante, Dr. Fernando.

Quero passar, então, para as suas considerações nesta audiência, a palavra ao Dr. Lucas Borges de Carvalho, Gerente de Projetos da Autoridade Nacional de Proteção de Dados.

O SR. LUCAS BORGES DE CARVALHO (Para expor.) - Bom dia!

Inicialmente quero cumprimentar todos e todas e agradecer a oportunidade de falar aqui em nome da ANPD.

Esse é um tema prioritário para a ANPD, em que a gente vem trabalhando com muito cuidado, estudando e enfrentando em várias em várias frentes. Então, eu vou trazer um pouco aqui do que a gente vem fazendo e da visão que a gente vem produzindo e preparando. Uma parte já foi apresentada, um estudo que a gente fez. Estamos, ainda, trabalhando numa proposta um pouco mais concreta que vai ser apresentada também a esta Comissão muito em breve.

Então, indo direto ao ponto, a principal questão, talvez, do ponto de vista da regulação, que é a abordagem que eu vou trazer aqui como membro de um órgão regulador, é a questão desse equilíbrio, porque, na verdade, toda a discussão em torno da regulação de IA passa por buscar um equilíbrio, na verdade, entre proteção de direitos fundamentais e inovação.

Então, quando a gente fala em viés, fala em discriminação, fala em transparência, fala na questão dos impactos que os sistemas de inteligência artificial podem gerar sobre as pessoas, na verdade, estamos falando em tentar minimizar riscos, em tentar proteger direitos. Mas, ao mesmo tempo, a gente sabe que, como acabou de ser dito aqui em relação à saúde, por exemplo, os sistemas de inteligência artificial envolvem também o desenvolvimento econômico, envolvem a produção de novas tecnologias e, na verdade, estamos em uma corrida global, eu diria até, do ponto de vista de como esse valor vai ser distribuído globalmente, inclusive, esses novos valores, seja do ponto de vista econômico, seja do ponto de vista do interesse público que os sistemas de inteligência artificial podem gerar.

Eu, ontem mesmo, li uma reportagem na *Folha de S.Paulo*, falando de um estudante de uma universidade americana que utilizou um sistema de inteligência artificial para poder ter acesso e conseguir ler um pergaminho histórico de uma biblioteca lá na Itália.

Então, os usos de inteligência artificial são, eu diria até, imprevisíveis e, talvez, infinitos, por isso é tão importante a gente, ao mesmo tempo, quando fala em proteger direitos, também não esquecer e tentar equilibrar a regulação do ponto de vista da promoção da inovação.

Bom, então, um primeiro ponto que eu acho que é muito importante nessa discussão é justamente ressaltar e enfatizar que há uma forte conexão entre proteção de dados pessoais e inteligência artificial, porque o treinamento de algoritmos pressupõe a utilização de uma quantidade massiva de dados e, em muitos casos, estamos falando de dados pessoais, principalmente naqueles casos mais controversos, que geram maior impacto, que são classificados como de alto risco.

Se a gente olhar o PL 2.338, as hipóteses previstas lá no PL como de alto risco, vocês vão ver que boa parte daqueles casos ali são casos em que estão envolvidos dados pessoais, como trabalho e emprego, controle de migração e fronteiras, área de saúde. Em muitas dessas situações, quando a gente fala, como eu falei, em discriminação, por exemplo, a gente vai falar, ou vai pressupor, ou estamos lidando com sistemas de inteligência artificial que lidam com dados pessoais. E a própria distinção entre dado pessoal e dado não pessoal é uma distinção que é muito fluida, porque mesmo em situações em que

estejam envolvidos dados não pessoais é possível ocorrer a inferência, é possível haver o cruzamento dessas informações com outras informações e dali serem revelados também dados pessoais.

E um terceiro ponto ali que eu trago, que é uma certa obviedade, mas precisa ser dita e enfatizada, é que a LGPD, assim como toda a legislação vigente, se aplica a sistemas de IA, mas a LGPD, em especial, porque, como eu falei, quando estamos falando de sistemas de IA, estamos também falando de proteção de dados pessoais. E reconhecer que a LGPD se aplica a sistemas de IA é importante, porque significa que tem uma série de pontos que precisam ainda ser esclarecidos em relação à aplicação da LGPD a dados pessoais. E esse é um movimento, eu diria, global, internacional, porque autoridades de proteção de dados de todo o mundo têm assumido um papel central na regulação de IA. Eu trago somente três exemplos recentes. Essa declaração conjunta das autoridades de proteção de dados do G7 sobre IA generativa, que é de junho de 2023, deste ano, justamente falando do ChatGPT, enfim, de todas essas IAs generativas, destacando ali as preocupações dessas autoridades em relação à fiscalização e ao monitoramento dessa nova tecnologia em relação tanto ao treinamento desses algoritmos com dados pessoais, como ao conteúdo que é produzido e à própria interação que há entre usuários e sistemas de IA generativa. Ali também há previsão de cooperação entre autoridades desses países desenvolvidos no acompanhamento dessa nova tecnologia, principalmente no impacto sobre a privacidade e a proteção de dados. A agência francesa CNIL também tem publicado, recentemente publicou uma nova consulta pública, tem publicado guias e trabalhado com *sandbox* regulatório na área de IA. E o ICO, que é a agência do Reino Unido, também tem definido claramente a IA como uma área prioritária devido ao potencial de alto risco sobre os indivíduos. Eu destaco um guia recente publicado pelo ICO, que é uma ferramenta de avaliação de riscos para desenvolvedores poderem desenvolver os seus sistemas com *Privacy by Design*, preocupados ali com o impacto sobre proteção de dados e privacidade.

Do ponto de vista da ANPD, é sempre importante frisar que a inteligência artificial é um ponto que está na agenda regulatória da ANPD. É claro que a gente vem acompanhando a discussão do Congresso Nacional, mas, como eu falei, a LGPD já se aplica aos sistemas de inteligência artificial que utilizam dados pessoais. Então, tem uma série de pontos que precisam ser esclarecidos. Não está definido ainda se será um regulamento, se será um guia, mas algumas questões, por exemplo, como a aplicação de princípios da LGPD; como o princípio da finalidade; o princípio da necessidade, que trata da minimização de dados; quais são as hipóteses legais que podem ser aplicadas para o treinamento de algoritmos, para a coleta de dados pessoais que vão ser utilizados para treinar sistemas de IA e, também, o famoso art. 20, da LGPD, que é o artigo que fala do direito à revisão de decisões tomadas com base em tratamentos automatizados de dados pessoais. Esse é um artigo também que traz uma série de questões que precisam ser definidas, incluindo a discussão sobre intervenção humana ou supervisão humana, que é um dos temas que está tratado aqui, com mais detalhes, claro, no PL 2.338. Isso está previsto na Fase 3, que é para o ano que vem.

Uma outra iniciativa recente da ANPD, também, é o estudo preliminar sobre o PL 2.338, que está disponível na nossa página na internet. Como eu falei, estamos trabalhando em uma segunda versão desse estudo, uma versão mais concreta, que vai ser entregue aqui à Comissão, inclusive com uma proposta de alteração do PL, uma proposta mais concreta. Mas as premissas são essas aqui, que esse próximo estudo, que vai ser em breve trazido aqui à Comissão...

Então, são dois pontos que eu destacaria. O primeiro é a questão da compatibilização da nova legislação com a LGPD. A gente sabe que o PL 2.338 tem um foco sobre a questão dos direitos e ele espelha, ele reflete, de algum modo, a própria estrutura da LGPD, seja do ponto de vista do foco na gestão de riscos, seja na existência, por exemplo, de comunicação de incidente de segurança, programas de governança, princípio da responsabilização e prestação de contas. E as próprias sanções administrativas que estão previstas no PL 2.338 são sanções muito similares àquelas sanções previstas na LGPD. Então, é preciso avaliar, de forma detalhada, o melhor mecanismo ou a melhor redação possível, para evitar que a nova legislação entre em conflito ou gere algum tipo de fragmentação regulatória, o que seria exatamente o contrário do que se busca, que é maior segurança jurídica, estabilidade e previsibilidade do ponto de vista da regulação. O segundo ponto, que é muito importante, é a defesa, como vários diretores da ANPD têm falado recentemente, do papel da ANPD como autoridade-chave no que se refere à regulação e à governança de IA no Brasil. O nosso entendimento é que a ANPD tem um papel central, dada essa forte conexão entre proteção de dados e inteligência artificial. O ideal é que a governança de IA esteja em harmonia, ou centralizada, ou associada com a regulação da proteção de dados pessoais.

Uma iniciativa recente, também relacionada com inteligência artificial, já com foco mais na inovação, é a consulta pública sobre *sandbox* regulatório, que foi aberta - que ainda está aberta, na verdade -, até 1º de novembro. O foco inicial deve ser na IA generativa. *Sandbox* é aquele mecanismo que permite trazer um aprendizado para o próprio regulador e, ao mesmo tempo, conferir segurança jurídica para os desenvolvedores que estão trabalhando com novas tecnologias, do ponto de vista da legislação de proteção de dados.

Este é o último eslaide. O tempo está se esgotando. A mensagem que eu queria deixar, em resumo, é de que uma regulação adequada de sistemas de IA pressupõe três pontos essenciais.

Primeiro, uma abordagem equilibrada entre direitos e inovação, seja o *sandbox* regulatório, seja a proteção de direitos e outras. A Estratégia Brasileira de Inteligência Artificial, que está prevista no próprio PL, precisa avançar, a nosso ver.

Segundo, é preciso harmonia e coerência com a legislação vigente, em especial a LGPD, a fim de evitar fragmentação regulatória e conflitos.

Em terceiro lugar, é preciso o reconhecimento do papel central da ANPD como órgão regulador, capaz de associar e trabalhar, de forma harmônica, a regulação de proteção de dados com regulação de IA, trazendo, claro, uma necessária discussão sobre o fortalecimento institucional da ANPD, porque a gente não pode trabalhar com a nova legislação sem pensar no fortalecimento dos órgãos reguladores que vão, na prática, implementar essa nova legislação.

Esses eram os pontos que eu queria trazer. Agradeço, mais uma vez, a oportunidade e me coloco à disposição para qualquer ponto.

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO) - Gostaria de agradecer ao Dr. Lucas, que vai acompanhar conosco as outras intervenções, as outras posições.

Gostaria também de registrar e colocar à disposição dos nossos expositores perguntas importantes de participação popular através sistema e-Cidadania. Vou registrá-las agora, e quem quiser pode, durante a sua fala, contemplar essas intervenções, essas solicitações e até comentários.

Bruno Ernesto, do Rio Grande do Norte, pergunta: "Quanto à automação processual, haverá espaço para decisões jurídicas por inteligência artificial?"

Helton Aparecido, de Minas Gerais: "Como será possível prevenir a veiculação de vídeos com conteúdo falso nas principais plataformas digitais [...]?"

Gabriel Pinho, do Rio de Janeiro, pergunta quais são as propostas e as limitações quanto ao uso da inteligência artificial em processos judiciais.

Ivan Sakima, de São Paulo: "Até que [ponto] [...] detalhes do direito material podem ser identificados pela IA no julgamento das ações?"

Luiz Cláudio, do Rio de Janeiro: "A inteligência artificial, quando bem utilizada, auxilia os diversos setores da sociedade, proibir [...] seria um erro [...]?"

Carla Alessandra, de São Paulo: "As questões éticas relativas à utilização da inteligência artificial devem ser tratadas com muita atenção, pois ainda é um campo de incerteza".

Fernando Mazzotta, de São Paulo: "Qual a principal garantia oferecida aos trabalhadores, sobretudo aos que possuem pouco acesso à tecnologia, [diante da] [...] substituição de [empregos formais] [...] por inteligência artificial?"

Luan Alves, do Rio de Janeiro, pergunta se é possível supervisionar o uso da IA de modo a evitar seu emprego para disseminação de *fake news*.

Letícia Garcia, do Rio Grande do Sul: "Está previsto que o Poder Judiciário faça aferição, por *software* ou algoritmos, da utilização abusiva da IA nas decisões judiciais?"

José Barbosa, de Goiás: "Qual a perspectiva de responsabilização por fatos originados direta ou indiretamente de algoritmos de inteligência artificial?"

Álvaro Cardoso, de São Paulo: "Qual a relação direta dos impactos do uso da inteligência artificial no contexto da economia e da produção industrial?"

Ederson Padovani, de São Paulo: "Profissionais de ciência de dados podem impactar a sociedade com suas soluções [...] uma regulamentação da área é necessária?"

Sebastião Aparecido, de Minas Gerais: "A inteligência artificial na esfera jurídica deve-se limitar à atividade meio, ou seja, atos de mero expediente".

Então, são essas as perguntas e os comentários. Peço a todos que queiram fazer...

Neste momento, passo a palavra ao Senador Marcos Pontes, Vice-Presidente da Comissão e também Ministro de Ciência e Tecnologia, especialista no tema.

Passo a palavra a V. Exa.

O SR. ASTRONAUTA MARCOS PONTES (Bloco Parlamentar Vanguarda/PL - SP. Para expor.) - Obrigado, Presidente.

Eu gostaria de só fazer alguns... Eu vou ter que sair para outra Comissão. O dia é meio corrido aqui, mas eu queria aproveitar este momento para, primeiro, parabenizar pela audiência. Realmente, são importantes esses debates. Nós temos

tido vários debates, de várias perspectivas, sobre a inteligência artificial, que, certamente, são importantes para compor a melhor solução para que essa lei possa ser abrangente o suficiente para proteger as pessoas, mas não restritiva de forma a não impedir o desenvolvimento tanto econômico quanto tecnológico do país através desse setor ou para esse setor. Então, são extremamente importantes essas conversas, ouvir os especialistas e aqueles que trabalham nas diversas áreas.

Como já foi falado aqui também, quando se trata de inteligência artificial, embora o pessoal pense só numa caixa-preta ali, que aprenda, mas todo esse processo de aprendizado de máquina quanto à utilização disso depende de dados, muitos dados, e aí vem o ponto realmente da parte de proteção de dados, porque esses dados, em grande parte, são relativos às pessoas, aí, entram as questões de ética, tudo que se envolve na proteção das pessoas.

E aí a Lei Geral de Proteção de Dados entra diretamente. Eu fiquei muito feliz em ver ali essa citação com relação a trabalhar alinhado, para que essas duas legislações trabalhem de forma alinhada. Não adianta a gente repetir, na legislação de inteligência artificial, algo que já está previsto na proteção de dados e muito menos contradizer, uma coisa que vá de encontro com o que já está escrito lá. Ou seja, precisa ser ao encontro, as coisas têm que funcionar corretamente.

Então, essa é uma coisa que tem que ser considerada, sem dúvida nenhuma. Eu tenho certeza de que o trabalho em conjunto com a Agência de Proteção de Dados vai permitir que essa legislação fique mais eficiente, a mais simples possível e de forma como não cause problemas. Tem que trazer soluções, e não causar problemas.

Acho muito adequado isso. Eu vejo nas perguntas a preocupação da população também com relação à utilização da inteligência artificial.

Uma coisa que eu já repeti aqui e é muito importante ter em mente: a inteligência artificial tem uma capacidade muito grande de resolver problemas, de mastigar - vamos chamar assim - uma quantidade muito grande de dados, e, a partir dali, chegar a alguma conclusão; mas ela não pode tomar decisões que vão afetar a vida das pessoas. Um ser humano precisa estar no *loop* para que se utilize a inteligência artificial como auxiliar, para que o ser humano tome a melhor decisão; mesmo porque é natural, na vida da gente - a gente nota isso todo dia, em tudo que a gente faz -, as decisões são, em grande parte, emocionais, e elas precisam ter esse componente emocional para que elas sejam mais justas, falando do ponto de vista de convivência, na parte interpessoal.

Então, parabéns pela apresentação.

Infelizmente, vou ter que sair, mas, Presidente, parabéns pela condução.

De novo, conte sempre com a gente.

Obrigado.

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO) - Muito obrigado, Senador Ministro Marcos Pontes.

Passo a palavra, neste momento, ao Diogo Cortiz, Coordenador do Mestrado e Doutorado em Tecnologias da Inteligência e Design Digital da Pontifícia Universidade Católica de São Paulo (PUC-São Paulo). Está conosco aqui, na videoconferência.

V. Sa. tem a palavra por até dez minutos.

O SR. DIOGO CORTIZ (Para expor. *Por videoconferência.*) - Obrigado, Senador Eduardo Gomes.

Cumprimento o senhor, cumprimento também todos da mesa. É um prazer estar aqui, discutindo esse tema tão importante, porque acho que a área de inteligência artificial realmente vai transformar todos os outros segmentos da sociedade.

Eu preparei alguns aspectos aqui para discutir, mas eu gostaria de dar um passo atrás e pensar: a gente está discutindo esse processo de regulação da inteligência artificial, e eu acho que é importante a gente retomar isso com mais força olhando para a Estratégia de Inteligência Artificial no Brasil. A gente está falando de uma regulação, sendo que, neste momento, a gente não tem uma estratégia, ou, na verdade, a gente tem uma estratégia, mas que não é tão bem detalhada, tão bem desenvolvida, sobre o ecossistema de desenvolvimento da tecnologia ou das tecnologias de inteligência artificial no Brasil.

Então, eu acho que um ponto importante que a gente precisa ter uma atenção maior é que a gente pode discutir regulação da inteligência artificial, mas também pode ser importante discutir regulação para a inteligência artificial; e aí, nesse sentido, entram em questão esses aspectos de pensar realmente uma estratégia de inteligência artificial.

No Brasil, existe a Ebia (Estratégia Brasileira de IA), que foi desenhada no ano passado, mas que, na minha visão, é totalmente insuficiente para o cenário brasileiro. Ela não aprofunda, ela não tem orçamento, ela não tem metas bem definidas. E uma coisa muito importante: a gente está entrando no momento mais oportuno das inteligências artificiais, que são esses modelos de linguagem e, quando a gente olha para a Estratégia Brasileira de Inteligência Artificial, em nenhum momento aparece, por exemplo, a palavra "português" ou a língua portuguesa, essa valorização, esse ativo importante

que o Brasil pode ter como um diferencial no desenvolvimento da tecnologia da inteligência artificial, que é justamente a incorporação, o desenvolvimento de modelos, de conjunto de dados, olhando para a língua portuguesa e também para a cultura brasileira.

Eu trago dois exemplos de fóruns em que isso foi discutido. Recentemente, o Comitê Gestor da Internet no Brasil organizou o Fórum Lusófono da Governança da Internet, que aconteceu no Museu da Língua Portuguesa em São Paulo, e eu fiquei responsável por moderar a mesa de inteligência artificial. Participaram representantes de indústria, da academia, de diferentes países lusófonos e se levantou bastante essa questão sobre a importância da criação de conjuntos de dados mais organizados sobre o português e a importância de os países que falam português estarem conduzindo esse processo até um ponto estratégico, porque hoje a gente tem esses grandes modelos de linguagem que trabalham com a língua portuguesa, mas a gente sabe que isso vai até um certo ponto. Eles têm alguma representação da língua portuguesa ali dentro, só que, quando você precisa de aspectos muito específicos da língua ou da cultura, esses modelos não são tão efetivos assim.

E, na semana passada, teve o Fórum de Governança da Internet, que é o Internet Governance Forum, que é um evento anual da ONU que aconteceu no Japão, e eu fiquei responsável também por organizar e moderar uma mesa discutindo o impacto da IA generativa no ecossistema da *web* como um todo. E participaram algumas *big techs*, como a Meta, a IBM e também alguns outros pesquisadores. Um dado muito importante foi apresentado por uma pesquisadora bastante conhecida na área de inteligência artificial, a Profa. Emily Bender, da Universidade de Washington, sobre o uso de dados e como esses modelos de linguagem são treinados predominantemente no inglês, muito poucas amostras de outras línguas, e, muitas vezes, você tem um processo de tradução ou de uma influência mesmo, tanto do ponto de vista da estrutura linguística como cultural dentro dessas respostas.

Então, eu estou dizendo isso porque acho que é importante a gente pensar no processo do desenvolvimento da inteligência artificial também no Brasil, porque senão, a gente vai ficar falando de regulação e boa parte dos sistemas que a gente vê - é óbvio que tem sistemas que são desenvolvidos no Brasil, por empresas nacionais, que são aplicados localmente -, você também depende de grandes provedores de tecnologia, no caso as *big techs*.

E aí entra uma questão que é essa falta de transparência que existe nesse tipo de tecnologia. A gente não sabe com quais dados eles foram treinados, a gente não tem informação sobre os modelos.

Inclusive, ontem saiu um artigo bastante interessante que cria um índice de transparência desses modelos de linguagem. Ele pega os dez principais, e nenhum consegue atingir a nota seis ali. São mais de cem indicadores que eles utilizam, desde o processo de treinamento dos modelos, o uso de dados, o que eles fazem para verificar o uso e as consequências. É interessante que até os modelos que são *open-source*, que são código aberto, como o Llama 2, que é um modelo que foi criado e desenvolvido pela Meta e é liberado de forma aberta, não conseguem atingir uma nota nesse sentido, porque falta transparência.

Recentemente, a gente vê esse movimento de criação de modelos *open-source*, por exemplo, o caso da Meta, só que eles apresentam como modelo *open-source* - aberto, que você pode usar, utilizar de maneira mais efetiva para o seu caso de uso ali - só que, ao mesmo tempo, falta transparência, porque a Meta - e recentemente eu estive na Meta na Califórnia - falou, por exemplo: "Para treinar o modelo de imagem, a gente utilizou todas as imagens que a gente tem públicas do Instagram", por exemplo. E eu disse: "Tá, mas tem mais coisa aí. Quais outros dados vocês utilizaram?", e eles não passam essa informação. Então, falta transparência nesse processo.

E, quando a gente está falando, então, de criar uma regulação ou políticas públicas para esse setor da inteligência artificial, eu acho que a gente precisa ter evidências, para ter uma coisa bem desenhada. Para algumas coisas, a gente tem evidências. A gente sabe que os modelos têm vieses, que eles podem prejudicar grupos específicos da população; disso a gente já tem evidências. Mas, para uma série de outras coisas, a gente tem um componente que é a imprevisibilidade: as coisas estão se desenvolvendo de uma forma muito rápida e cada vez a gente vai dar um uso específico para a tecnologia, principalmente por conta desses modelos fundacionais. E esses modelos fundacionais a gente pode entender como um sinônimo de IA generativa, porque a gente consegue desenvolver uma série de outras aplicações sobre esses modelos, e isso traz uma complexidade muito grande. Essa falta de transparência eu acho que é um fator que deveria ser um foco muito importante dentro de um ambiente de regulação, que deveria ser o passo 1 para a gente começar a entender melhor sobre esse processo.

Tem algumas outras questões aqui de coisas práticas desse cenário. Eu trabalho bastante com pesquisas nesse ecossistema digital, principalmente olhando para a *web*, para a internet e a *web* de forma geral, e justamente, nesse painel que a gente organizou durante o Internet Governance Forum, nesse evento da ONU, a gente focou muito no impacto, por exemplo, do uso da inteligência artificial dentro desses vários serviços da internet, da *web*, principalmente no impacto econômico para o sul global. Um dos pontos que a gente levantou, que eu deixo também para depois possivelmente ser aprofundado, é o uso mais excessivo dessas IA generativas em serviços de busca, por exemplo. Então, quando você vai em um buscador hoje, você faz uma consulta, ele vai gerar uma lista de *links*, e, por exemplo, se você perguntou algo que tem no meu *site*,

ele vai dar o *link* para o meu *site*, você vai acessar o meu *site*, você vai conhecer mais sobre os meus serviços e, se eu tiver algum tipo de anúncio, eu vou ser remunerado por isso.

Hoje, com o uso, a adoção desses sistemas dentro da internet, você acaba mudando toda a dinâmica dessa economia digital, o que é algo também com que a gente tem que estar atento, porque tanto o Google está fazendo movimento nesse sentido, como outras empresas, como a Meta, estão oferecendo vários serviços para mudar exatamente todo esse ecossistema digital. Então, essa é uma questão que a gente deveria olhar.

Eu sei que está acabando o meu tempo e eu gostaria de agradecer.

Muito obrigado.

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO) - Eu quero agradecer a exposição feita agora, com bastante competência, pelo Diogo Cortiz.

Vamos só fazer uma alteração aqui na ordem, uma coisa rápida, até por uma questão de preferência nominal.

Nós vamos ouvir agora a Dra. Nina da Hora, Diretora-Executiva do Instituto Da Hora, que também está presente por videoconferência.

V. Exa. tem o prazo de até dez minutos.

A SRA. NINA DA HORA (Para expor. *Por videoconferência.*) - Bom dia a todos. Obrigado, Senador, pelo convite. Bom dia, colegas. É um prazer mais uma vez estar com alguns de vocês aqui, que têm participado desde o início da implementação da Comissão, para discutir a regulação de IA e também para dar início a outras discussões antes dessa proposta de regulação, que tangencia como que a gente pode comunicar melhor, para a sociedade, o que quer dizer IA e quais as diferenças das ferramentas que derivam desse conceito e das suas construções, assim como os diversos seminários e encontros nas universidades e fora delas, para tentar abranger e disseminar de uma forma mais qualitativa.

E assim eu inicio a minha fala. Eu acredito que tem a sua importância a população estar participando mais de audiências públicas como esta, acho que nós tivemos também um grande ganho em relação ao projeto de lei das *fake news*, um grande ganho da mídia e da população que passou a perceber e a querer participar das discussões.

Mas, para além das colocações que foram feitas aqui pelos colegas, eu queria trazer um tema relacionado à educação de como que infelizmente nós, às vezes, replicamos opiniões e visões que não são visões da realidade da construção e do desenvolvimento de IA no Brasil. Acho que vale a lembrança - apesar de que eu não sou uma pessoa que tenha nascido na década de 80 - de que o Brasil tentou... Muitos pesquisadores que ainda estão vivos, as universidades como a USP, a Unicamp, a PUC-Rio, universidades pioneiras na discussão de como seria possível implementar e pensar tecnologias de IA a partir do contexto brasileiro.

Então, vale resgatar que isso foi pensado, tiveram tentativas, inclusive tentativas em parcerias dessas universidades que eu citei com os governos da época, e implementações de comissões, como esta que foi implementada no Senado, que infelizmente não foram para frente. Então, historicamente, o Brasil não deveria ser colocado como um país que não tem pessoas capazes, pesquisas e infraestrutura necessária para fazer essa discussão de IA. Se a gente for resgatar a história, teve momentos, tanto políticos como também de interesses individuais, empresariais ou não, que acabaram desvirtuando o caminho que esses pesquisadores estavam tentando construir, inclusive, na figura, na representação dessas universidades que eu trouxe aqui.

Hoje, nós nos encontramos no que eu chamo de um campo minado de como que a gente consegue se relacionar bem com o posicionamento nacional de você desenvolver tecnologias e a indústria nacional a partir de pesquisas, sejam pesquisas acadêmicas ou pesquisas que são direcionadas pela indústria nacional. Também nós temos a dificuldade de conseguir relacionar isso com as desigualdades que aumentaram no Brasil.

No início da audiência, o doutor colocou muito bem as alternativas, as propostas que podem beneficiar uma área da medicina, e isso só é possível de ser aferido e compartilhado como ele fez quando existe investimento nas pesquisas, sejam pesquisas aqui ou pesquisas fora. Então, eu queria colocar minha preocupação de que, na discussão de regulação da qual eu tenho participado desde o início da discussão do projeto de lei com a Comissão de Juristas... Eu acredito que tenhamos que trazer mais perspectivas acadêmicas de diferentes áreas da inteligência artificial para fazer essa discussão de uma forma um pouco mais qualitativa. Então, de fato, eu sinto falta de a gente valorizar o que foi construído e o que é construído no Brasil. Eu acho isso importante, porque nós temos um contexto bem peculiar em relação aos outros países com os quais nós temos parceria, com os quais nós estamos fazendo essa discussão, e essas peculiaridades precisam ser levadas em consideração. Eu vi o colega falando da língua portuguesa e eu trabalho mais com imagem.

Então, um ponto que tem sido crítico, que tem sido parte da discussão de regulação e com o qual eu fico feliz é como que nós podemos trabalhar os vieses da perspectiva da ética, não só os vieses raciais, como também os vieses de gênero. E existem propostas que estão sendo pensadas, que são propostas multidisciplinares.

Então, quando eu falo da perspectiva da academia, eu não estou falando só da computação, que é a minha área; eu estou falando da antropologia, eu estou falando da história, eu estou falando das ciências sociais, eu estou falando para além da matemática e da computação, que, muitas vezes, trabalham com certezas, só olhando para os números. Nós precisamos ter esses estudos comportamentais, esses estudos de como que os vieses estão impactando a nossa sociedade de uma forma qualitativa até para apresentar para o Senado, para contribuir para uma regulação para que ela não caia nessa... Para mim, eu considero uma desinformação o que tem sido colocado de que a proposta da regulação pode paralisar o desenvolvimento científico e tecnológico, o que não é uma verdade. Nós temos experiência o suficiente de que nós precisamos alinhar desenvolvimento técnico e científico com os direitos do cidadão. Então, eu considero que regulações como essas fazem parte da construção e da continuidade de uma verdadeira democracia.

Obrigada.

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO) - Eu quero agradecer a exposição.

Continuam chegando perguntas do e-Cidadania e, daqui a pouco, eu vou passar outras perguntas.

Neste momento, nós vamos ouvir o Dr. Marcelo Finger, Professor Titular do Departamento de Ciência da Computação do Instituto de Matemática e Estatística da Universidade de São Paulo (USP), que nos acompanha por videoconferência. V. Sa. tem o prazo de até dez minutos para a sua exposição. Bom dia.

O SR. MARCELO FINGER (Para expor. *Por videoconferência.*) - Bom dia, Senador. Bom dia a todos presentes aqui. Agradeço muito a oportunidade de estar falando com todos aqui. Sou professor da área de Ciência da Computação e trabalho com pesquisas de inteligência artificial há mais de 30 anos. Além disso, falo normalmente com a população em uma coluna de rádio sobre temas relativos à inteligência artificial. Gostaria de trazer alguns pensamentos para cá e talvez até responder a algumas das perguntas que foram feitas pelos ouvintes que se manifestaram.

Em primeiro lugar, com relação a proibir a inteligência artificial, gostaria de trazer o caso da Itália, que proibiu o uso do programa chamado ChatGPT e foi um tremendo fracasso. Não funcionou proibir; já reverteram isso. Então, as iniciativas que visam meramente proibir o uso da inteligência artificial, de certa forma, não estão resolvendo o problema, porque a tecnologia continua se desenvolvendo e, draconianamente, simplesmente falar que é proibido não funciona.

Além disso, gostaria de trazer ao foco aqui uma experiência recente e bastante bem-sucedida, no meu ponto de vista, que foi o tratamento da inteligência artificial na negociação entre os estúdios de Hollywood e os escritores e roteiristas que estavam em greve. Uma primeira coisa a que gostaria de chamar a atenção é que, pela primeira vez, numa negociação salarial ou de condições de trabalho, a tecnologia foi posta em primeiro plano, e muitos dos acordos firmados têm a ver com o uso da inteligência artificial. Então, basicamente, isso parece indicar que a tecnologia, hoje em dia, é uma condição de trabalho e pode ser tratada como tal.

Em segundo lugar, gostaria de chamar a atenção de que, nessa negociação, a inteligência artificial não foi proibida, não foi proibida nem pelo lado dos estúdios, nem pelo lado dos escritores. Ambos estão autorizados a utilizá-la, com algumas restrições. Então deixe-me explicar o que foi: há um *framework*, há um arcabouço legal, em torno do qual esses acordos foram firmados, segundo o qual é o seguinte: os estúdios gostam de manter o direito autoral sobre os roteiros e os *scripts* escritos, e, também, dentro desse arcabouço legal, textos produzidos por inteligência artificial não são passíveis de direito autoral. De certa forma, os estúdios concordam que os escritores precisam escrever, e não a inteligência artificial gerar uma versão, porque se for uma versão totalmente gerada pela inteligência artificial, ela não é passível de direito autoral e não é passível, portanto, de ter a posse ou a transferência. Isso é o arcabouço legal.

E com o que eles concordaram? Eles concordaram com o seguinte: se um estúdio quiser gerar por inteligência artificial um esboço de um roteiro, esse roteiro é transferido para os escritores para eles desenvolverem a partir daí. Só que isso não é uma eliminação de trabalho, ele é apenas a utilização da inteligência artificial para a agilização do processo, porque os escritores recebem pelo texto produzido pela inteligência artificial como se eles mesmos tivessem escrito. Então, não há supressão de ganhos ao usar a inteligência artificial.

Por outro lado, os escritores podem usar a inteligência artificial para gerar trechos e adaptar partes. Eles não estão proibidos de utilizar a inteligência artificial na relação dos textos deles. Eles também têm como agilizar a produção dos textos utilizando as ferramentas.

Por fim, os estúdios não podem usar roteiros antigos que têm em posse como base para gerar novos roteiros. Isso ocorre por vários motivos, um dos quais é que ao colocar um roteiro como *prompt* ou como algo num sistema, essa informação

passa a ser pública e pode ser utilizada pela empresa que detém os direitos do *software*. Por exemplo, se colocassem no ChatGPT, a OpenAI poderia usar esse texto que contém os direitos, e os direitos autorais estariam quase que perdidos dessa forma. Então, esse foi um acordo muito interessante entre dois lados contenciosos, em que ninguém está proibido de utilizar, a utilização foi regulamentada, e isso foi algo completamente novo.

Isso também é muito específico porque havia um *framework* legal no caso dos direitos autorais sobre textos e roteiros, que, por exemplo, já não está tão claro assim na manipulação de imagens.

Não sei se vocês sabem, a greve dos roteiristas terminou; a dos autores, acho que não. E aí o problema é enorme porque há a manipulação de *software* e a manipulação de imagem, uma vez que você pode escanear o seu rosto e falar alguma coisa, e o *software* já utiliza suas imagens e sua voz para aparecer você falando num vídeo outras coisas, coisas que você nunca falou na vida. Isso é muito chocante. A qualidade ainda não é maravilhosa, mas é impressionante. Então, a possibilidade de gerar notícias falsas que possam invadir a internet é muito grande: vídeos, fotos e textos falsos.

As empresas de IA estão muito preocupadas com isso. Eu acho que deveria haver um esforço maior de inserção de marcas d'água em imagens e vídeos gerados pela empresa para que qualquer um possa verificar: "Ah, esse vídeo foi realizado por tal programa"; "Ah, esse vídeo, eu não sei qual é a origem dele, eu não vou mexer". Então, precisa de um desenvolvimento tecnológico aqui também para facilitar o uso.

A lei tem uma série de direitos de imagem que não estão claros. O ator, por exemplo, pode fazer um filme em inglês e aí aparece ele mexendo a boca e falando em português. Será que ele tem que receber por essa modificação, ou o estúdio tem direito de fazer isso? Esse tipo de contencioso não está resolvido.

Então, tem uma série de pontos que ainda precisam ser melhorados, talvez com base em algum *framework* legal que estabeleça esses pontos.

Os estúdios, no caso da greve dos roteiristas, aceitaram pagar pelo roteiro escrito pela inteligência artificial. Não acho que há uma compreensão nesse sentido no caso das imagens.

Outra coisa importante que muito poucas pessoas falam é sobre a utilização da voz. Hoje em dia, eu trabalho nessa área de pesquisas de diagnóstico de doenças respiratórias pela voz. Será que é lícito, daqui a 20 anos, pegar uma imagem, por exemplo, essa aqui e falar: "Ah, Marcelo, você estava sofrendo de tal e tal doença naquela época, consigo detectar pela sua voz". Isso deveria ser permitido, se eu não autorizasse e se eu não estiver mais vivo? É uma série de pontos que a gente precisa levar em consideração na análise da voz. Ele é muito bom, a gente pode, através da voz, com um celular, ajudar a estabelecer o estado de saúde de uma pessoa em qualquer lugar, pode ser na enfermaria do Hospital das Clínicas, em São Paulo, ou pode ser numa barca no Rio Tapajós, com uma facilidade enorme, mas a gente não deveria permitir que certas coisas fossem usadas fora desse contexto de tratamento.

Então, isso é também um ponto que eu gostaria de levar e, portanto, já vou concluir falando que, basicamente, a minha mensagem é: proibir não adianta. Olhem para os casos interessantes em que houve resolução de assuntos em que não proibiram o uso da IA e permitiram às partes usar.

Existem dificuldades em áreas que não têm um *framework* legal para serem usadas, e não ignorem novas possíveis tecnologias, como as tecnologias que utilizam a voz para fazer detecção do estado de saúde das pessoas.

E com isso, termino a minha participação.

Agradeço aqui, olha, bem na música. Muito obrigado a todos.

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO) - Muito obrigado, Dr. Marcelo Finger.

Passo, neste momento, a palavra ao Dr. André Carlos Ponce de Leon Ferreira de Carvalho, Diretor do Instituto de Ciências Matemáticas e de Computação da Universidade de São Paulo (ICMC-USP), representante de Thais Vasconcelos Batista, Presidente da Sociedade Brasileira de Computação (SBC), que está também conosco em videoconferência para a sua intervenção, pelo prazo de até dez minutos.

Tem a palavra V. Sa.

O SR. ANDRÉ CARLOS PONCE DE LEON FERREIRA DE CARVALHO (Para expor. *Por videoconferência.*) - Senador Eduardo Gomes, demais autoridades, colegas, bom dia. Agradeço pela oportunidade.

Então, o que eu vou falar um pouco vai ser abrindo uma outra vertente com relação à regulação, que são as oportunidades que nós temos de ter cuidado para não inibir que ocorra no Brasil. Vou falar um pouco sobre o panorama do mercado mundial e como está a parte de ensino e de talentos em IA no mundo inteiro.

Tem uma estimativa de que em 2032 o mercado de inteligência artificial vai chegar a US\$2,5 trilhões no mundo. E o Brasil tem todas as condições, talentos nós temos, nós temos uma juventude muito habilidosa que pode participar dessas iniciativas.

Esse outro quadro mostra a participação dos diferentes continentes ou regiões do mundo nesse mercado de IA em 2022. A menor participação é da América Latina e do Oriente Médio, que juntos dão 14,26% do mercado de inteligência artificial. Esse outro quadro mostra o investimento: que país está investindo e quanto está investindo em IA. Destacam-se os Estados Unidos, não só no período que vai de 2003 a 2022, mas também em 2022, seguido pela China e Reino Unido. O Brasil não aparece nesse quadro.

Esse outro quadro mostra os investimentos que os países têm em investimento privado em IA. Mais uma vez, o Brasil não aparece, mas na América Latina aparece a Argentina, que em 2022 foi um dos países que mais investiu de forma privada na inteligência artificial.

Todo esse investimento tem criado uma grande demanda e procura por talentos em IA. O mundo está contratando gente talentosa em IA. Não são poucos os casos de talentos que perdemos do Brasil para o exterior. Não só jovens formandos, como também docentes de universidades brasileiras. Isso é preocupante, porque nós vamos perder uma massa crítica muito relevante para que o país tenha um protagonismo na inteligência artificial.

Este quadro mostra também a procura por empregos em IA no mundo, que tem disparado, e mostra o quanto nós formamos em inteligência artificial em relação aos países do mundo. Ou seja, o Brasil até que é um dos que forma mais, mas está na rabeira, está um pouco abaixo da Holanda, mas a Holanda tem menos de 10% da população brasileira. Ou seja, nós formamos muito pouca gente, e isso acaba restringindo também as nossas oportunidades.

Em 2018, o relatório da mais importante universidade chinesa na área de tecnologia, a Tsinghua, avaliou o panorama da IA no mundo, principalmente em relação aos talentos em IA. Identificou que no mundo existiam, em 2017, cerca de 200 mil talentos em IA - mais de 10% deles estavam nos Estados Unidos -, e em segundo lugar vinha a China. Ela também avaliou onde estão os talentos do topo, os melhores talentos em IA. Mais uma vez, os Estados Unidos ficaram em primeiro lugar, com 5,5 mil talentos, num universo menor, e a China caiu para o oitavo lugar.

Este quadro mostra a proporção entre talentos que estão no topo de IA e os talentos de IA. Os Estados Unidos tinham 18,1%, a China, 5,4%, mostrando a deficiência que eles têm de talentos no topo. O Brasil tem 6,5%, mas isso é por quê? Porque a gente tem poucos talentos. Portanto, os poucos que nós temos no topo acabam contribuindo para aumentar esse número.

Neste outro quadro, temos o número de talentos topo e a proporção de talentos topo em relação a todos os talentos. Novamente, os Estados Unidos ficaram em primeiro lugar - como eu já mostrei -, e a China em sexto lugar, enquanto o Brasil não aparece.

Aqui, temos a distribuição dos talentos de IA nas universidades do mundo inteiro. Basicamente, na América do Norte, na Europa, muitos na Índia, e alguns na América Latina, somente no Brasil.

E chegamos às universidades que têm mais talentos em IA no mundo. Em quantidade de talentos em IA, a China é a maior universidade, em Tsinghua, com 822; depois vem China, Índia, China, Estados Unidos, China, China, China. Isso em talentos em IA. Quando chegamos aos talentos topo em IA, a distribuição modifica um pouco: o Brasil e a América Latina continuam tendo poucos, mas quando nós olhamos quais universidades têm mais talentos no topo em IA, em 2017, a Universidade de São Paulo aparecia em quinto lugar. Mas como uma andorinha só não faz verão, 60 pessoas numa universidade não vão fazer com que o Brasil tenha um protagonismo em IA no mundo.

Tem outro estudo interessante, que é da Fapesp (Fundação de Amparo à Pesquisa do Estado de São Paulo), que mostra a publicação em IA. A publicação tem relação forte com talentos topo em IA. Então, o Brasil estava atrás do Canadá - estava bem atrás -, e estava um pouco acima da Coreia do Sul. E, quando nós pegamos a lupa e olhamos como o Brasil trabalha em IA e qual a produção brasileira em IA, a USP aparece - também a USP é a maior universidade, tem vários grupos de pesquisa, mais que outras universidades também, ela aparece no topo -; tem a Unicamp; UFPE; UFMG; a Unesp também, boa parte das universidades públicas do país.

Tem outro quadro mais recente, que é onde o Brasil está no panorama global de IA, segundo o índice de IA. Nesse índice - no caso, aqui, estão os Estados Unidos -, ele avalia o talento, infraestrutura, o ambiente operacional, as pesquisas, o desenvolvimento, a estratégia governamental em IA, a parte comercial, a escala e a intensidade. Então, temos Estados Unidos, China, Singapura... Continuando, vimos que o Brasil está na 35ª posição.

Por outro lado, quando avaliamos a primeira coluna - que é essa que eu marquei com o círculo -, o Brasil está em 21º lugar, justamente em talentos de IA. Ou seja, a gente está melhor em talento do que na nossa situação global em IA. Mas isso é bom? Não, porque a posição nº 21 ainda é pouco, diante do que o Brasil conseguiria fazer.

Então, se olharmos apenas o talento, vamos ver que o Brasil está à frente de várias nações desenvolvidas, muitas da Europa também, mas, mesmo assim, nós estamos muito... Estamos acima da Áustria, por exemplo; acima da Noruega; acima da Itália, mas é muito pouco diante do que o Brasil consegue e poderia ter.

Então, o que eu queria falar? Eu queria primeiro mostrar uma oportunidade que o Brasil tem e algumas preocupações. Nós precisamos de uma regulação que evite abusos, mas temos que aproveitar também essa oportunidade para criar oportunidades para o país também, para o Brasil desenvolver a IA e ter um protagonismo mundial. Nós vimos que temos talentos, mas poucos; temos *startups* em IA também no Brasil, mas também são poucas; e temos alguns centros de pesquisa, que são poucos e trabalham com poucos recursos.

E se esse panorama não mudar, nós vamos perder o bonde da história. Então, vamos ter uma ótima regulação de IA, mas não vamos ter um protagonismo internacional ou uma IA forte em relação ao que a gente espera do Brasil.

Muito obrigado mais uma vez pela oportunidade.

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO) - Muito obrigado. Impressionante a apresentação do ponto de vista de necessidade de investimento e oportunidade para abrigar e formar talentos em inteligência artificial - um dado importante.

Eu gostaria só de - vamos passar para as duas últimas exposições - fazer uma observação para que fique registrada na nossa audiência pública. Com o ajuste da aprovação pela Comissão, pelo Presidente Carlos Viana e pelo Vice-Presidente, Senador Marcos Pontes, e a dinâmica, nós tomamos a decisão, com a anuência do Presidente Rodrigo Pacheco, de, além das audiências públicas, todas que serão formadas, ao final nós termos uma Comissão Geral, que é o debate no Plenário do Senado Federal.

Mas eu gostaria também, já que é a primeira vez que fazemos essas audiências públicas e precisamos dar publicidade, além do fato de serem todas registradas e estarem à disposição da população brasileira no sistema da internet, pela competência, capacidade de capilaridade, junto às assembleias legislativas, câmaras municipais, o sistema de informação do Interlegis, gostaria de pedir à TV Câmara, em nome do Diretor Érico Silveira e a todo o seu corpo editorial - eu tenho certeza de que o assunto é prioritário, é um assunto que permeia todos os outros assuntos -, formalmente, solicitar à linha editorial da TV Senado que avalie a possibilidade de formularmos algum programa fazendo a síntese das audiências públicas para que a gente possa, através de todos que participaram das audiências públicas, divulgar o trabalho que está sendo feito, justamente no sentido de colher contribuições e dar ciência.

E falo isso porque se tem uma conclusão a que o Relator já chegou, através da conversa com a Comissão de Juristas, com os Parlamentares, com os setores de controle, de organização, as plataformas é que nós teremos que ter uma lei viva, uma lei que consiga regular sem impedir.

Então, quanto mais repercussão, quanto mais condições nós pudermos fazer... e se a gente pede isso aos outros quanto mais a um veículo de comunicação da Casa. Então, a competente TV Senado, eu realmente acho que se for enquadrar na programação a síntese dessas audiências, facilmente nós podemos, pelo menos para a população da TV aberta, ter cravado o lugar de consulta, que é através do e-Cidadania, mas que as pessoas possam ver e tirar suas dúvidas e as suas contribuições através da comunicação do Senado.

Então, fica aí o meu pedido. Gosto muito do nosso diretor, tenho conversado com ele sobre outros projetos, mas peço especialmente que a TV Senado dê uma atenção melhor para produzir um material que possa ser mais compreensível pelas populações, por todos, já que aqui existem muitos setores representados.

Então, passo a palavra, neste momento, ao Dr. André Lucas Fernandes, diretor e fundador do Instituto de Pesquisa em Direito e Tecnologia do Recife (IP.rec), que também nos acompanha por videoconferência.

V. Sa. tem o tempo de até 10 minutos para a sua exposição.

O SR. ANDRÉ LUCAS FERNANDES (Para expor. *Por videoconferência.*) - Bom dia a todos e a todas presentes.

Eu gostaria de cumprimentar os Exmos. Senadores componentes desta Comissão Temporária Interna sobre Inteligência Artificial e o faço na pessoa do Sr. Senador Relator Eduardo Gomes.

Eu gostaria ainda de agradecer a oportunidade de contribuir com uma discussão tão importante como esta, sendo um advogado, pesquisador e ativista dos direitos digitais do Nordeste brasileiro. Eu falo em nome do IP.rec (Instituto de Pesquisa em Direito e Tecnologia do Recife), organização que faz parte da Coalizão Direitos na Rede e é um centro

independente e multidisciplinar de estudos sobre direito, tecnologia e sociedade, e que tem por missão a produção de pesquisas, capacitação e comunicação sobre os temas que lhe dão nome.

Temos como orientação de atuação o compromisso fundamental com a democracia, a construção multissetorial da ciência e da regulação e a participação pública significativa nos temas de interesse da sociedade.

Senadores e Senadoras, se a esta Casa cabe, dentro das competências constitucionais, debruçar-se sobre este tema de relevante interesse, é preciso, então, situá-lo na perspectiva de uma paisagem sociotécnica e jurídica em sentido estrito, nesta zona de comunicação entre o sistema do direito e o sistema econômico. Existe uma falsa percepção de que o que está em jogo aqui e que convoca o Congresso Nacional a um esforço regulatório é o uso de sistemas de IA para questões que não envolvem diretamente pessoas, seres humanos.

Defender que nosso debate é isso, especialmente atrelado a um subtema de regulação como algo para emperrar a inovação e desenvolvimento, é defender fumaça e opacidade. É impreciso afirmar que uma legislação com bases gerais de adequação e garantia de direitos e mapeamento de riscos seja algo desconhecido do setor privado ou do setor público, por exemplo.

Ao menos duas razões me chegam neste momento. Primeiro, que a existência de uma regulação detida comparada ao cenário de vácuo legislativo só aumenta o nível de estabilidade, previsibilidade e segurança para inovar. Segundo, que tal tipo de informação ignora o processo de boas práticas no desenvolvimento tecnológico, que envolve aquilo que, no jargão técnico, nós chamamos de "documentação" e cuja prática ajuda em perspectivas de inovação colaborativa e aberta a otimizar soluções, corrigir erros e gerar riqueza com menor custo.

Documentar, ou seja, registrar as etapas de desenvolvimento, dos erros, das melhorias, desde a ideia até o lançamento do produto no mercado é uma boa prática base no setor.

O Prof. Luciano Floridi, que é um especialista no assunto e veio aqui ao Congresso Nacional ano passado, tem um artigo no qual ele aborda de forma taxativa que o tempo da autorregulação no setor de tecnologias digitais acabou. Argumentos sobre a falta de capacidade cognitiva do Estado, ou seja, que o Estado não entende de tecnologia para poder regular são falsos na raiz, afinal, nunca se tratou na história de um domínio onisciente do Estado sobre todos os temas que são postos na mesa e são de interesse social.

O equívoco se desdobra na falsa investida contra uma regulação dita, entre aspas, "geral". A generalidade apenas poderia ser acolhida para alguns como principiologia, o que é que isso quer dizer. Ou então nós teríamos regulações setoriais apenas e que não seriam a agenda imediata deste Poder Legislativo.

Se o conceito tratado de IA não é unívoco, esse é um problema deixado para a ciência e não para a produção legislativa. É cristalino que a definição legal que é proposta no texto do 2.338, por exemplo, tem caráter operativo e, ademais, que o "geral" aqui não se opõe a "setorial". A norma jurídica, por outro lado, é um instrumento de ordenação e incentiva as práticas sociais positivas. Atingir, portanto, o núcleo dos riscos e desmontar uma falácia desregulatória está ligado à criação de um conjunto de obrigações em sentido estrito que permita construir um cenário de transparência, explicabilidade e esclarecimento sobre como essas tecnologias funcionam, do dado de entrada ao produto de saída. Em suma, tudo o que envolve o tema relativo da justiça algorítmica.

"Lixo entra, lixo sai" é uma regra famosa na ciência dos dados. Cenários fantasiosos que ignoram a realidade demográfica brasileira e prometem um futuro ideal com base numa receita um tanto cansada do norte global não deveriam nos interessar. A realidade da desigualdade, que atinge quase 50% da população do Brasil, sendo a maioria negra, e os expõe à adoção massiva de sistemas de vigilância, reconhecimento facial e falsos positivos nos interessa como premissa legislativa.

Por outro lado, focar nas previsões de fenômenos climáticos com a IA, e usando a IA como sendo algo mais importante do que lançar luzes sobre práticas ambientais exploratórias que um modelo de extração inadequada de dados e da natureza pode causar no planeta e no Brasil e que cria um cenário irreparável e insustentável de escassez, envenenamento ambiental e morte, não nos interessa.

O aspecto setorial e o incentivo à produção destas tecnologias cabe ao Poder Executivo, como foi mencionado aqui. Nesse ponto, a Estratégia Brasileira de Inteligência Artificial (EbIA), já mencionada, é uma das poucas iniciativas concretizantes envolvendo a criação, por exemplo, de laboratórios temáticos na própria EBS é uma das poucas iniciativas - e mesmo essa iniciativa está envolta com alguns limitadores, como foi mencionado, como recursos irrisórios, falta de transparência na destinação desses recursos e opacidade sobre alguns parceiros privados que estão envolvidos no processo.

É na competência executiva, portanto, das agências já existentes que esse recorte deveria ser tratado. Ao Poder Legislativo cabe o papel de dialogar com duas grandes necessidades focadas na produção normativa: a necessidade do contexto tecnológico, que está sendo posta aqui em alguns momentos, e as necessidades geradas pelas consequências desse contexto, que são as necessidades sociais, públicas ou privadas.

Dentro desses pontos equívocos que eu mencionei antes, eu gostaria de destacar algumas das boas e más propostas do PL 2.338, inclusive atendendo a um chamado do Relator que vi no primeiro dia sobre dialogar diretamente com o texto. Parece-me que a parte dos fundamentos e princípios colocados no texto são adequados e dialogam com as técnicas de interpretação jurídica mínimas e mais atuais do debate científico sobre a IA e suas consequências. O conceito do art. 4º, entretanto, enfrenta o desafio da amplitude do tema, do termo inteligência artificial na forma de sistema de inteligência artificial. Notem, Exmos. Senadores, que o termo IA foi pensado como uma estratégia de *marketing* para a captação de investimentos científicos lá na origem nos Estados Unidos. Não há no nome "inteligência artificial" um valor científico.

O problema maior, contudo, no texto, está na dicotomia conceitual entre fornecedores e operadores de sistema de IA. Aqui, sim, juridicamente falando, a gente tem consequências práticas relevantes, porque esses atores de fornecimento e operação terão atrelados a si regimes de responsabilidade, tanto em nível de *compliance* como em nível de danos causados. E a dicotomia, que serve muito bem ao regime do Código de Defesa do Consumidor, a LGPD, não revela clara a complexidade no desenvolvimento da IA e de seus múltiplos usos.

A gente precisa ter um tratamento granular desses atores ou a criação de hipóteses que diferenciem quando o "autor", entre aspas, da IA não está envolvido em um contexto de desenvolvimento para o mercado, o que é uma diferença extremamente relevante e que fala sobre a realidade do desenvolvimento da IA por pessoas independentes.

Um elemento ajuda a diminuir esse problema no texto, entretanto. Quando a gente faz uma interpretação sistemática da parte de riscos, como eles estão propostos, eles estão colocados de forma granular, escalonada e pensados, como já foi mencionado aqui antes, como um mecanismo atualizável. Ou seja, há espaço para a gente melhorar isso ao longo do tempo.

É importante mencionar, reforçando uma mensagem contra pretensas brigas de torcida, que também já foi mencionado aqui sobre PLs, que existe uma defasagem epistemológica científica que torna os textos de outros projetos de lei, em que pese o valor histórico que eles tenham, datados e, portanto, imprecisáveis ao propósito regulatório.

O acúmulo diário de estudos e análises e críticas e os avanços de 2019 para cá depõem contra essa anacronia do texto anterior e dos outros textos propostos e, ao mesmo tempo, reforça a necessidade de chamar um marcador jurídico sobre esse tema. O sistema jurídico recorta aquilo que é, para ele, Srs. Senadores, relevante. Essa relevância é política e social. Nesse recorte, quando falamos de direitos das pessoas ou direitos humanos, é porque uma marcação foi feita sobre um fato bruto, tornando-o um fato jurídico e gerando efeitos. Esses efeitos são os direitos que a gente menciona quando fala de direito à privacidade ou dos direitos bem elencados no art. 5º do PL 2.338.

Então, quando existir um sistema de IA em operação e esse sistema tocar na esfera de uma pessoa, devemos ter um processo regulatório de IA incidente sobre o tema. Essa seleção, como já dito, é operacional e serve ao desenvolvimento de toda a sociedade.

Um outro ponto importante - já encaminhando mais para o final da minha fala - é a questão do *compliance* que é proposta nos arts 19 e 26. Na verdade, é a questão da responsabilidade civil que é proposta logo após esse regime de *compliance*. O texto merece melhorias nesse ponto. A responsabilidade civil se situa na articulação de uma última fronteira de compensações quando a gente tem alguma consequência relevante e negativa. Só que, se no art. 27 parece estar descrito lá a responsabilidade objetiva, que seria a mais adequada para o caso concreto, as exceções do art. 28 podem criar um cenário de completo esvaziamento prático do instituto.

A gente está falando aqui de *enforcement* e da necessidade de atacar como esses produtos tecnológicos e oriundos da tecnologia e da economia da informação têm sido desenvolvidos de maneira impermeável ao escrutínio público. E essa opacidade é resultado de um setor desregulado e permeado por práticas exploratórias.

A produção de indicadores, por exemplo, sobre gado brasileiro e a predição de modelos climáticos não são o tema em questão. Os temas que nos importam são os modelos que vão criar ranqueamento de pessoas, vão excluir ou incluir políticas públicas ou, então, que vão, de maneira arbitrária, gerar concessão de crédito ou precificação de produtos e perfisamentos de toda a ordem, envolvendo, inclusive, crianças e adolescentes no ambiente escolar.

Enfim, para terminar, eu acho também que vale citar, quanto ao aspecto do fomento e da inovação, o texto é bastante tímido. Do art. 38 em diante, a gente basicamente fica com a menção a *sandbox* regulatórios. Só que o Brasil já tem um marco legal de ciência, tecnologia e inovação.

Também o ponto relativo a direitos autorais é só um piso que está posto no debate do 2.338 e a gente precisa avançar muito mais nisso. O tema e o texto poderiam trazer, concordando com outros colegas - e finalizando -, proposta sobre letramento para o desenvolvimento e uso da tecnologia, porque não avança, de um lado, com a necessidade brasileira nem, de outro, com os melhores padrões de participação pública significativa e educação.

É importante dizer, Excelências, que o Brasil tem um compromisso para se desenvolver tecnologicamente com preservação ambiental. Quando a gente fala então sobre esses desenvolvimentos mais opacos e que não levam em conta o custo também ambiental dessa tecnologia, a gente não está falando de uma verdadeira inovação, que é aquela que está preconizada na Constituição Federal de 1988. A gente pode, sim, pensar numa inovação atual, não dogmática, participativa e aberta, que tenha foco social e que permita que o Brasil mostre ao mundo o desenvolvimento da IA para o novo e não para essas receitas cansadas que já estão aí.

Eu agradeço.

O tempo é bastante curto para poder tentar dar conta de todo o processo, mas coloco o trabalho do IP.rec, do qual faço parte, à disposição.

E desejo que os trabalhos da Comissão sejam os melhores possíveis.

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO) - Agradeço o Dr. André pela intervenção.

Fico muito feliz porque essas audiências públicas têm essa visão - e é essa realmente a grande vantagem do Parlamento: os relatórios vêm a Comissões específicas, temáticas ou temporárias, para os melhoramentos. E os melhoramentos só acontecerão com sugestões, sejam elas pertinentes ou não, atuais ou não, importantes ou não. O importante é que elas ocorram, por isso que há uma série de audiências públicas no processo de construção de convergência, espelhando a realidade. Afinal de contas, constitucionalmente é responsabilidade do Congresso, com o amparo do conhecimento, prover legislação.

Passo a palavra agora, antes, porém, fazendo uma grata observação que do meu estado, que está aqui presente o Deputado Estadual Gutierrez Torquato, Deputado Estadual da cidade de Gurupi, e também o Deputado Eduardo Fortes - dois grandes amigos, jovens Deputados Estaduais do nosso jovem Estado de Tocantins. São momentos dignos de lembrança do nosso saudoso Siqueira Campos ter aqui no Senado tocantinenses, Senadores e Deputados estaduais. Então sejam bem-vindos a essa Comissão Especial de Análise da Regulação da Inteligência Artificial no país, a qual preside o Senador Carlos Viana e eu tenho a responsabilidade de ser o Relator dessa matéria, que vai modificar a vida do cidadão e a política na Câmara dos Deputados, no Senado, nas Assembleias Legislativas, nas Câmaras de Vereadores, enfim, da população brasileira. É um prazer tê-los aqui. É uma honra.

Passo a palavra nesse momento a Gustavo Zaniboni, Presidente da Coordenação de Inteligência Artificial da Associação Brasileira de Governança Pública de Dados Pessoais (govDados), que vai fazer a sua exposição, que é a última dessa audiência pública, mas eu continuo dizendo a todos que, mesmo através do Sistema e-Cidadania e das redes de comunicação do Congresso Nacional e do Senado Federal, estão disponíveis todas as exposições feitas nessas e nas outras audiências de instrução dessa importante matéria.

Tem V. Sa. o prazo de até dez minutos para a sua exposição - se precisar, um pouco mais.

O SR. GUSTAVO ZANIBONI (Para expor.) - Eu agradeço ao Senador Eduardo Gomes. Agradeço a todas as pessoas que estão nos acompanhando.

Eu venho representando a govDados para tentar trazer uma visão técnica para tentar apoiar e ajudar nesse processo de regulação legislativa. Vou tentar fazer algumas observações que me são caras e que entendo que precisam ser melhor discutidas para que o resultado do processo legislativo seja aplicável na prática.

Eu começo a apresentação só lembrando aos senhores que tudo o que puder ser automatizado vai ser automatizado. E isso vale desde sempre. O ser humano tenta automatizar tudo. Agora, nós estamos vendo uma mudança enorme, que já vem acontecendo nesse século, mas, de 2022 para cá, todo mundo passou a entender que existe alguma coisa que se chama inteligência artificial e que isso está revolucionando a forma de automatizar tarefas, principalmente porque o público em geral passou a ter acesso a esse tipo de sistema. Antigamente, isso estava muito dentro das empresas, estava muito dentro do meio acadêmico, e, com o ChatGPT, com o MidJourney, todo mundo passou a ter acesso a esse tipo de sistema. E agora a gente tem que lidar com esse tipo de novidade tecnológica. Por isso, é importante a gente olhar um pouquinho com cuidado alguns pontos.

Primeiro, a gente está num processo legislativo que já está em andamento - teve um relatório de quase mil páginas, que serviu de base para este projeto de lei -, mas eu quero trazer a importância da precisão técnica para a regulamentação. E aqui, surrupiando do Fernando Pessoa - "navegar é preciso" -, estou falando de precisão nos dois sentidos: de ser preciso a gente fazer essa abordagem técnica, e de que essa abordagem técnica tem que ser precisa, no sentido de a gente ser muito objetivo e ter muita certeza do que a gente está falando. Eu falo isso, porque a gente tem uma dificuldade já no começo da definição de inteligência artificial. E, se a gente for olhar o que está escrito lá no PL, ele dispõe sobre o uso da

inteligência artificial, só que a gente está tentando regular a tecnologia, a inteligência artificial, ou a gente está tentando regular o uso da inteligência artificial? Então, a primeira dúvida, a primeira questão que a gente tem é como a gente vai abordar esse tipo de questão.

A minha proposta é para a gente começar a olhar um pouco e entender que inteligência artificial tem a ver com como as coisas em computação são feitas e não com o que é feita. E é importante a gente ter esse tipo de sentimento para a gente poder olhar e entender o que tem risco, o que causa dano e o que a gente precisa regular.

Eu vou dar um exemplo muito rápido. Imaginem uma seleção automatizada de candidatos.

Eu posso simplesmente ter um sistema computacional que faz uma série de perguntas para o candidato, que pega o currículo dele, que pesquisa algumas palavras-chave nesse currículo dele e que decide se esse candidato vai ou não para a próxima etapa do processo seletivo. Então, eu estou fazendo uma seleção automatizada. Eu não estou usando inteligência artificial nesse processo que eu descrevi para vocês.

Eu posso ter o mesmo processo que faz algumas perguntas para o candidato, só que, em vez de eu fazer um *search* de algumas palavras no currículo dele, eu aplico um sistema computacional que vai procurar contexto e que vai tentar interpretar o currículo dele, mas, ainda assim, eu faço um *score*, eu faço uma conta, uma média, um cálculo e defino, a partir desse cálculo, se ele vai ou não para a próxima etapa. Eu já estou usando um pouco de inteligência artificial nessa avaliação contextual do currículo dele, mas, ainda assim, a tomada de decisão está sendo simplesmente um *score*. Não tem inteligência artificial na tomada de decisão.

Ou eu posso ter o mesmo sistema com o mesmo resultado, com o mesmo objetivo, em que ele não dá uma nota para cada resposta e procura certas palavras no currículo do candidato, mas ele faz uma avaliação - e o programador não sabe muito bem como é que ela é feita, porque ela é baseada em dados que são utilizados para treinar um modelo de inteligência artificial - e ela me dá a resposta se aquele candidato deve ou não prosseguir para a próxima fase do processo seletivo.

Assim, eu tenho o mesmo "o que" - eu estou fazendo uma seleção automatizada de candidatos -, só que eu tenho três "comos": eu tenho um como que não tem nada a ver com inteligência artificial; um como que tem um pedaço que usa inteligência artificial, mas que não é o decisório; e um como que é baseado em inteligência artificial para tomar uma decisão.

Deixe-me ver se eu passo ali. Não, não está passando... Ah, desculpe. Foi mal. Então, vamos lá.

O primeiro ponto que a gente tem que discutir, que a gente deveria olhar com um pouco mais de cuidado, é a própria definição de inteligência artificial que está na lei, porque, se essa definição não for muito benfeita e ela não puder ser aplicada na prática, o escopo da lei vai passar a ser as empresas tentando se enquadrar ou se desenquadrar na lei, olhando a definição de inteligência artificial. E, quando a gente vê o PL e a gente vê as coisas que estão descritas no PL, eu fico com a sensação de que nós estamos falando de sistemas computacionais de apoio à tomada de decisão ou de sistemas computacionais de tomada de decisão, que podem ou não usar inteligência artificial. Apesar de, no texto do PL, a gente estar falando em inteligência artificial durante todo o texto ou de sistemas de inteligência artificial, quando a gente descreve, quando a gente vê as questões que estão colocadas no PL, a gente não está falando da tecnologia de inteligência artificial; a gente está falando de sistemas que eventualmente podem ou não usar a inteligência artificial. E daí a gente fica com aquela pergunta: se usar a inteligência artificial, eu tenho que fazer tudo o que está na regulamentação e, se não usar, eu não preciso fazer para fazer a mesma coisa? Então, esse é um ponto importante que a gente tem que considerar.

E é muito importante que, quando eu uso inteligência artificial para resolver algum problema, eu tenho a imprevisibilidade. Quando eu não uso, quando eu tenho um sistema computacional normal, o programador explica certinho para o computador o que é para ele fazer. Quando eu uso inteligência artificial, é o próprio sistema que, a partir dos dados, entende e interpreta como ele tem que resolver: ele vai buscar resolver e dar a saída mais próxima possível dos dados que foram apresentados a ele. E aí eu tenho o componente da imprevisibilidade.

Outro ponto do PL importante para a gente olhar é o ponto que fala de transparência, explicabilidade, inteligibilidade e auditabilidade. O que é possível, quando eu tenho um sistema que usa inteligência artificial, é falar em transparência e auditabilidade. Explicabilidade do modelo de inteligência artificial que está ali dentro e que usa inteligência artificial e inteligibilidade eu não consigo fazer se tem inteligência artificial. Se eu pudesse explicar como um modelo toma uma decisão ou como ele pega aquela entrada e chega àquela saída, eu não ia usar inteligência artificial, eu ia fazer um programa normal de computação, que pegaria aquela entrada e levaria àquela saída. Se eu sou capaz de explicar por que aquela saída aconteceu, eu não preciso usar inteligência artificial, que é mais caro computacionalmente, que é mais difícil de se implementar e que é imprevisível. Então, aqui a gente tem no texto: "explicitando a lógica e os critérios relevantes para a produção de resultados". Eu consigo explicar como o sistema funciona - eu consigo ter transparência do sistema

e auditabilidade -, mas eu não consigo explicar exatamente por que ele deu aquela resposta ou por que ele errou aquela resposta.

Eu vou dar um exemplo rápido para vocês no próximo eslaide, mas esse é o paradoxo da caixa-preta. O modelo não é opaco, não é caixa-preta, porque as pessoas impedem que você olhe o sistema. Nesses sistemas novos do ChatGPT, os modelos têm bilhões de parâmetros. Mesmo que eu abra um modelo para que a gente olhe dentro desse modelo, mostre o algoritmo, mostre os dados, mostre o modelo, é impossível saber por que ele está tomando aquela decisão. E a resposta geral para qualquer sistema de inteligência artificial do porquê ele tomou aquela decisão eu explico: ele tomou aquela decisão, porque ela é a mais próxima estatisticamente dos dados que foram usados no treinamento daquele modelo, daquela rede neural.

Nesse aspecto, como eu consigo fazer só transparência e auditabilidade e eu vou ter que pensar o que é explicabilidade - a explicabilidade do sistema e não do modelo em si -, eu preciso ter uma importância muito grande, dar um peso muito grande na governança dos sistemas que usam IA.

Eu vou avançar, só para mostrar o exemplo.

O exemplo que eu trouxe mostra... Esses resultados são resultados errados. Aquele número seis ali de cima ou aquele número cinco ali foi uma rede neural simples, muito simples, que pega um caractere desenhado à mão e transforma num número. A gente não consegue explicar por que ele não definiu aquele primeiro número como o número seis. Simplesmente não sei explicar. Ele acertou 97,7% das vezes um número escrito à mão, alguns muito mais complicados do que aquele número seis ali do primeiro quadradinho, mas aquele ele errou. E eu simplesmente, como desenvolvedor, não tenho como explicar por que ele errou.

Por isso, o que a gente tem que fazer... E até o Senador Marcos Pontes podia estar aqui para eu comentar. A sugestão é que a gente use um processo muito parecido com o que se faz na aviação. Na aviação, eu tenho uma regulamentação, eu tenho um processo muito rigoroso de desenvolvimento de sistemas e de desenvolvimento de *software* para aviação, mas cada avião é diferente, cada tipo de arquitetura é diferente. É mais ou menos o que acontece quando temos modelos de inteligência artificial. E onde que eu pego? Eu pego no processo. Eu garanto que o processo é o melhor processo possível. Eu garanto que eu tenho auditabilidade e que eu tenho transparência no processo. Se o avião cair e eu tive essa transparência, eu não vou ser culpado pela queda do avião. O FAA vai chegar lá ou a Anac vai chegar lá e vai entender o que tem que mudar no processo; vai-se melhorar o processo para que os próximos desenvolvedores não derrubem o seu avião.

Eu falo também de viés de discriminação e ética. Rapidamente, tem todos os aspectos de viés de discriminação que a gente ouve falar que os sistemas podem ter por dados errados, por não estar fazendo uma avaliação correta, por não estar treinando com toda a população que tem que ser treinada, mas tem um tipo de viés que é muito importante. Inteligência artificial repete o passado. Se eu tenho um processo, por exemplo, de uma seleção para o curso de enfermagem - e você tem a maior participação de mulheres no curso de enfermagem, historicamente -, quando eu for automatizar esse processo usando um algoritmo que use inteligência artificial, muito provavelmente ele vai repetir o passado, assim como qualquer sistema que usa inteligência artificial faz. Se eu tenho esse tipo de comportamento, eu preciso eticamente entender se faz sentido eu automatizar esse processo, porque eu não vou conseguir mudar o passado.

Aí entra uma parte do comitê de ética em que eu olho: eticamente, faz sentido eu automatizar um processo que já é discriminatório ou que já tem um viés? E não é um viés dos dados no sentido de eu estar com os dados errados, os dados estão corretos; é o viés da sociedade. A gente tem que levar em consideração que inteligência artificial repete o que está nos dados. Se o nosso processo como seres humanos é enviesado, tem discriminação, ao automatizar esse processo, eu vou trazer isso para dentro do sistema. Então, não é só uma discriminação algorítmica; é uma discriminação do mundo que está sendo levada para dentro do sistema. E eu tenho que ter como impedir que isso aconteça.

Mais uma questão dos pontos específicos do PL: os agentes de tratamento. Eu ter só dois agentes de tratamento não reflete a realidade. Eu tenho: quem desenvolveu o algoritmo, *transformers*; quem treinou o algoritmo; quem fez o modelo; quem refinou o algoritmo; quem fez o sistema que faz a engenharia para injetar os dados de *input* nesse algoritmo; e um aplicativo que usa um ou mais desses algoritmos, no final das contas. Eu não tenho dúvida de que tem, sim, um fornecedor ou um fabricante que empacota aquele sistema. Lembrem-se de que falei de sistema automatizado e de tomada de decisão? Quem faz o sistema, certamente, tem que ter responsabilidade, mas eu não posso simplesmente achar que eu só tenho um operador e um fabricante. Eu tenho toda uma cadeia, inclusive, os *foundations models*, que são os modelos de base, que são modelos básicos em cima de quem você trabalha. Eles são como se fossem um alicerce. A gente está ouvindo muito falar disso nos modelos de linguagem. Tem muitos produtores de *foundations models* que estão preocupados, porque eles estão fazendo um modelo genérico que tem várias camadas antes de isso ser usado de fato em alguma aplicação. E é aplicação. É o sistema que pode causar um dano que pode causar um problema. Então, a gente tem que conseguir, pelo menos, separar isso para que a gente dê tranquilidade para o desenvolvimento dos modelos dos *foundations models*.

Tem todo o problema de treinamento, de propriedade intelectual e uso de dados pessoais. Aqui, a gente tem que seguir o que a LGPD exige.

E a gente tem que só tomar um pouquinho de cuidado, principalmente, com os modelos de linguagem. Tem muita gente falando: "Não, vocês estão usando dados que têm propriedade intelectual".

A primeira etapa de um modelo de linguagem - ChatGPT - desse tipo é uma "tokenização", que pode ser considerada praticamente uma "anonimização". Você não vai achar dados pessoais nem dados com propriedade intelectual dentro do modelo.

Mas olhem que interessante, olhem a "matemática" funcionando. Esse mesmo modelo em que você não consegue encontrar, lá dentro, esses dados - esses dados não existem lá dentro - consegue produzir um texto que pode ter dados pessoais e pode ter matéria, material e informação com propriedade intelectual.

Então, a gente tem que olhar como abordar isso. Isso não está muito abordado no PL.

A LGPD fala de tratamento automatizado que envolva dados pessoais. Então, a gente tem que fazer aqui uma conversa com a outra legislação e tentar só entender os casos particulares quando a gente tem sistemas que usam inteligência artificial.

Bom, daí, finalmente, tem toda a parte da abordagem de risco. O maior problema dessa abordagem de risco é quando a gente já tem o risco *a priori* definido.

Por exemplo, eu sei que um sistema de biometria é considerado risco alto. Daí eu pergunto: a biometria que destrava o meu celular, para fazer com que eu acesse o meu conteúdo aqui no meu celular, tem alto risco? Ou é só a biometria que ajuda a identificar um cidadão para fazer um *score* social, como acontece na China por exemplo?

De novo, eu estou apontando o dedo para uma tecnologia - biometria - quando o problema está no uso dessa tecnologia, está lá no sistema.

Então, eu deveria ter um sistema de alto risco ou até de risco extremo em que não possa ser implementado o monitoramento social das pessoas com câmeras e a biometria facial em espaços abertos; eu tenho um intermediário em que eu procuro pessoas ali para segurança pública, que tem que ter todos os cuidados para que o modelo não tenha o viés, não repita uma condição de discriminação, até tente melhorar essa condição; e eu tenho um sistema que destrava o meu celular, que tem zero risco. E tudo é biometria!

Então, se eu vou com abordagem de risco *a priori*, falando que biometria tem risco alto, eu estou trazendo todas essas coisas. Eu deveria olhar não para a tecnologia, mas para o uso da tecnologia.

E como colocar isso no regulamento, na lei é muito importante, porque, dependendo de como eu colocar, ou eu estou pesando a mão para aplicação que não tem risco de fato, com esse risco *a priori*, ou eu estou tirando o pé quando eu tenho o risco muito alto.

Risco é o sistema, é o aplicativo. Não é a tecnologia.

Agora, quando eu uso inteligência artificial, eu tenho alguns aspectos adicionais para levar em conta. Quais? Toda a parte de vies de discriminação de que eu falei e toda a parte que eu falei, lá no começo, sobre a imprevisibilidade.

Olhem que interessante: se eu sou uma empresa e quero fazer uma avaliação, em vez de usar a inteligência artificial, eu posso usar um método que tem um risco muito maior para o cidadão que não use inteligência artificial e que não vai estar sob o escrutínio da lei. Se eu usar inteligência artificial, que, teoricamente, melhora a situação para o cidadão, eu estou com todo esse escrutínio.

Talvez, nisso, a gente tenha que tomar um pouco mais de cuidado.

(Soa a campanha.)

O SR. GUSTAVO ZANIBONI - Já encerrando, eu tenho que lembrar se eu tenho impacto do algoritmo ou impacto do aplicativo.

Eu quero só responder porque daí eu já respondo a uma das perguntas que foram recorrentes ali.

O que aconteceu, além de todo esse problema da regulamentação que a gente tem dos sistemas que usam inteligência artificial? Eu tenho uma outra parte: o mundo mudou. Antigamente, para eu fazer um comercial como o da Elis Regina ou para eu pegar uma foto do Donald Trump sendo preso, falsa, *fake*, e conseguir produzir esse material, eu tinha que pegar um profissional, com Photoshop, com aplicativos de imagem, para produzir essa foto. Demorava. Hoje, qualquer pessoa consegue escrever um *prompt* num aplicativo e ter essa imagem em alta resolução muito fácil.

Eu não vou resolver esse problema regulando o uso de inteligência artificial. Não adianta pôr marca d'água. Não adianta. A gente tem que entender que o mundo mudou e que nós vamos ter problemas que, antes, a gente não tinha. Antes de ter

carro, a gente não tinha problema de acidente de trânsito. Depois que a gente teve o carro, a gente passou a ter problema de acidente de trânsito.

Então, a gente tem que conseguir diferenciar o que a gente vai regular.

Aqui, regulando tecnologia ou mesmo regulando sistemas, eu não consigo evitar que usos, principalmente, de IA generativa comecem a trazer alguns problemas, problema de propriedade intelectual, de plágio, problemas de golpes.

Tem gente dando golpe no WhatsApp em que pega um pedaço da voz da pessoa e finge que é o filho pedindo dinheiro para a mãe, com a voz da pessoa. Isso, a gente não vai resolver regulando a inteligência artificial. Isso, a gente vai ter que entender que o nosso mundo mudou e que a gente precisa olhar para essas coisas também.

Então, eu quero deixar aqui, de novo, só reforçando, que ter uma precisão técnica, ter o alicerce bem-feito ajuda no processo legislativo, para que os juristas consigam escrever muito bem a lei.

Eu fico à disposição e trago como sugestão que estes pontos sejam olhados com mais cuidado: definição, agentes e se a gente está falando mesmo de tecnologia ou se a gente está falando dos sistemas, porque eu tenho muito medo de a gente estar regulando uma tecnologia quando a gente quer regular um sistema.

Se a gente for regular o sistema, a gente tem que diferenciar quando eu tenho ou não a tecnologia que tem a imprevisibilidade, que são a maioria dos sistemas que usam inteligência artificial, que usam redes neurais.

Obrigado.

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO) - Quero agradecer ao Gustavo Zaniboni.

Nada mais havendo a tratar, declaro encerrada a presente sessão, agradecendo a participação de todos.

(Iniciada às 10 horas e 45 minutos, a reunião é encerrada às 12 horas e 31 minutos.)