



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

31/10/2023 - 12ª - Comissão Temporária Interna sobre Inteligência Artificial no Brasil

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP. Fala da Presidência.) - Bom dia a todos.

Havendo número suficiente, declaro aberta a 12ª Reunião da Comissão Temporária Interna sobre Inteligência Artificial no Brasil.

A Comissão foi 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 de substitutivo sobre inteligência artificial no Brasil, criada pelo Ato do Presidente do Senado nº 4, de 2022, bem como eventuais novos projetos que disciplinem a matéria.

Informo que esta reunião se destina à realização de audiência pública com o objetivo de debater a aplicação e implicações da inteligência artificial nas eleições e na disseminação de informações, os desafios e os riscos que a inteligência artificial apresenta à integridade jornalística e à democracia, em cumprimento ao Requerimento nº 4, de 2023, de autoria do Senador Eduardo Gomes.

O público interessado em participar desta audiência pública poderá enviar perguntas ou comentários pelo endereço [www.senado.leg.br/e-cidadania](http://www.senado.leg.br/e-cidadania) ou ligar para 0800 0612211.

Novamente, as pessoas interessadas em participar, enviar suas perguntas ou comentários participem pelo Portal do Senado, [www.senado.leg.br/e-cidadania](http://www.senado.leg.br/e-cidadania), ou liguem para 0800 0612211.

Encontram-se presentes no plenário da Comissão:

- a Sra. Tainá Aguiar Junquillo, Professora de Direito, Inovação e Tecnologia do Instituto Brasileiro de Ensino, Desenvolvimento e Pesquisa, o IDP - eu convido a Sra. Tainá, por favor, para ocupar o lugar junto conosco;
- a Sra. Patricia Blanco, Presidente do Instituto Palavra Aberta - Patricia, por favor;
- o Sr. Diogo Rais, Diretor-Geral do Instituto Liberdade Digital - por favor, Diogo, venha conosco.

Encontram-se também presentes, por meio do sistema de videoconferência:

- João Caldeira Brant Monteiro de Castro, Secretário de Políticas Digitais da Secretaria de Comunicação Social da Presidência da República - obrigado pela participação, João;

- Sra. Renata Mielli, Coordenadora do Comitê Gestor da Internet (CGI.br) - o.k., obrigado pela participação;
- Sra. Luciana Moherdau - desculpe se eu falei o nome incorretamente -, que é Pesquisadora do Grupo Jornalismo, Direito e Liberdade da Escola de Comunicação e Artes da Universidade de São Paulo - obrigado pela participação;
- Sra. Dora Kaufman, Professora da Pontifícia Universidade Católica de São Paulo, a PUC-SP - Dora, obrigado pela participação;
- Sra. Celina Bottino, Diretora de Projetos do Instituto de Tecnologia e Sociedade - Celina, obrigado pela participação também;
- Sr. Marcelo Antônio Rech, Presidente da Associação Nacional de Jornais - Marcelo, obrigado pela participação.

Deixe-me falar algumas palavras aqui sobre esta reunião: esta audiência faz parte de um conjunto de audiências que nós temos conduzido aqui no Senado com o objetivo de orientar esse marco da inteligência artificial.

Sem dúvida nenhuma, a inteligência artificial já está presente no nosso dia a dia e vai estar cada vez mais; tem sido um desenvolvimento de muitas décadas, embora nos últimos anos é que tenha sido mais comentado, falado a respeito de inteligência artificial, mas é um desenvolvimento de muitas décadas. E, como toda tecnologia disruptiva, ela tem possibilidades de ajudar, e muito, em muitos dos problemas que nós temos que resolver no nosso planeta, como mudanças climáticas, epidemias, pandemias, resolução de problemas imensos com relação à fome, à pobreza, às questões de água e preservação do que nós temos no nosso planeta, e também na melhoria dos trabalhos, dos sistemas de mais eficiência, tanto público quanto privado, ou seja, ela tem uma possibilidade gigantesca de ajudar em muitas áreas.

Mas, como toda tecnologia disruptiva, também tem as suas possibilidades de impactos negativos. Basta a gente pensar, por exemplo, em tecnologia nuclear, que certamente tem uma utilização muito positiva quando se trata de saúde, desenvolvimento de novos fármacos, diagnósticos, etc., mas que pode também ser levada para o lado errado, na produção de armamento nuclear e assim por diante. Ou seja, como toda tecnologia, nós precisamos garantir que aqui no Brasil nós tenhamos o desenvolvimento da pesquisa, o desenvolvimento do conhecimento sobre inteligência artificial, as tecnologias derivadas desse conhecimento, a aplicação dessas tecnologias nos diversos ramos de negócio, para que nós tenhamos desenvolvimento econômico e nós tenhamos competitividade no Brasil e não fiquemos para trás com relação aos outros países, porque eles vão continuar a desenvolver lá e nós não podemos ficar para trás, como nós já ficamos em muitas das tecnologias no passado. Mas, por outro lado, a gente tem que garantir também a liberdade das pessoas, a gente tem que garantir a ética, a não discriminação e muitos outros fatores que podem ser potencializados pelo uso de inteligência artificial - lógico que sempre como resultado da ação humana, que está sempre por trás de tudo isso.

Então, aí vem a importância de nós termos um marco legal que seja, ao mesmo tempo, flexível o suficiente para garantir o desenvolvimento da tecnologia, a utilização dessa tecnologia para a competitividade do país e para geração de empregos, ou para a transformação de empregos compatíveis com o futuro, que já está batendo à porta em muitos dos setores, mas também responsável o suficiente para proteger as pessoas. É esta a nossa intenção aqui hoje, falando a respeito de jornalismo, das notícias, de como elas saem, das eleições e outros temas extremamente importantes no nosso contexto do dia a dia. Faz parte, então, de um conjunto.

Nós já tratamos aqui das aplicações, aliás, nós estamos vendo essa tecnologia por quatro perspectivas: a primeira é o desenvolvimento da tecnologia; a segunda é a aplicação dessa tecnologia nos diversos setores; a terceira é a implicação dessa aplicação; e a quarta é o invólucro jurídico para que tudo isso seja feito da maneira correta e nós tenhamos uma lei que perdure, para que não se faça uma lei extremamente restritiva que, no dia seguinte, já esteja desatualizada, porque a tecnologia acelera muito rapidamente, mas também uma lei que proteja os cidadãos e que não entre em choque com outras leis já existentes. Lembro que existe isso, nós tivemos muito disso aqui, essas discussões, por exemplo, com a Lei de Proteção de Dados, que já existe, não faz sentido criar uma lei com as mesmas coisas que já existem na Lei de Proteção de Dados e outras leis setoriais, que já trabalham na regulação de cada setor, por exemplo, nas agências como Anvisa, Anatel, ANP e assim por diante. Então é importante sempre ter isso em mente. Às vezes a gente quer colocar tudo dentro de uma lei só, mas já existem outras, então não precisa fazer isso; seria até ineficiente e muitas vezes incorreto.

Como nós vamos funcionar aqui hoje?

Nós temos nove apresentadores, então nós vamos destinar um tempo de dez minutos para cada apresentador. Aqueles que estão aqui têm a facilidade de ver o relógio ali, conseguem ver o tempo facilmente. Faltando um minuto, vai tocar uma sirene como esta aqui, é fácil ouvir, e, faltando 15 segundos, vai entrar uma voz feminina muito convincente dizendo que faltam 15 segundos. Para aqueles que estão remotos, vocês não têm essas ferramentas aqui, então eu peço que controlem o tempo no relógio, no celular, de alguma forma, controlem o tempo, mas aos 15 segundos vocês vão ouvir a voz feminina falando a respeito dos 15 segundos finais. Lembro que isso é um tempo importante para que nós cumpramos nosso tempo total de duas horas, isso dá uma hora e meia de tempo corrido, fora este tempo que já estou usando, ou seja, a gente vai chegar ao finalzinho. Depois que todos falarem e nós recebermos as perguntas do e-Cidadania e também as perguntas locais, nós vamos abrir para as questões da seguinte forma: eu vou passar a palavra para cada um dos nossos apresentadores ou debatedores, para suas considerações finais e responder às perguntas também ao mesmo tempo. Dado o tempo restrito que nós temos, acho que vai funcionar bem, o que não impede de um apresentador fazer perguntas para o outro também poder responder, sem problema nenhum.

A nossa sequência de apresentação que recebi aqui será a seguinte: primeiro, Diogo Rais, depois Celina Bottino, depois Luciana Moherdau, depois Patricia Blanco, depois Marcelo Rech, depois Renata Mielli, Tainá Aguiar Junquilha, Dora Kaufman e João Brant, nessa sequência. Então, só para vocês ficarem atentos aí na sequência.

Sem mais, então, eu concedo a palavra para o Sr. Diogo Rais, que é Diretor do Instituto Liberdade Digital, para sua apresentação de dez minutos.

Obrigado, Diogo.

**O SR. DIOGO RAIS** (Para expor.) - Perfeito, Senador. Eu que agradeço. É uma honra enorme poder estar nesta Casa, em especial discutindo um tema de tamanha relevância, com tantos colegas que admiro tanto e com quem tenho aprendido tanto ao longo do tempo.

Então, eu queria cumprimentar a todos - todos que estão presentes, seja presencialmente, seja virtualmente -, afinal, uma das maravilhas da tecnologia é justamente o desafio de enfrentar as barreiras físicas do tempo e do espaço e colocarmos todos juntos, independentemente de partilharmos o mesmo espaço e até o mesmo tempo. Então, de uma certa maneira, todos estamos juntos aqui nessa missão.

Queria me apresentar brevemente, até para que eu possa ser honesto com meu ponto de partida. Meu nome é Diogo Rais, eu sou Professor de Direito Eleitoral e de Direito Digital da Universidade Presbiteriana Mackenzie, co-Fundador do Instituto Liberdade Digital, que é um *think tank*, um centro de pesquisa dedicado aos temas sobre tecnologia e democracia.

Já desde 2009 eu tenho me dedicado ao tema da internet nas eleições, basicamente como a tecnologia poderia impactar as eleições. Nos últimos anos, desde 2015, um tema central foi o tema da desinformação *online* ou das famosas *fake news*. Desde então, eu tenho tido a oportunidade de pesquisar e publicar a respeito desses temas sempre com um olhar nas eleições.

De uma certa forma, é um pouco esse o meu ponto de partida. E eu gostaria de trazer algumas considerações a respeito do enorme desafio que é colocado à nossa frente aqui e que se traz com enorme honra para que eu possa compartilhar e tentar colaborar - esse gigante desafio do nosso tempo e, de uma certa forma, um desafio também global.

Bom, a partir desse olhar das eleições e da desinformação, eu gostaria de trazer talvez a minha fala compartilhada em dois grandes eixos: um eixo específico sobre a atuação da inteligência artificial nesse cenário, um cenário de eleições e desinformação; e depois, além das preocupações, algumas sugestões para que a gente possa olhar para isso com o olhar possível que o direito possa empregar a essa matéria.

Bom, de uma certa forma, o direito, assim como todas as demais ciências ou técnicas, tem os seus limites, e muitas vezes a gente se esquece um pouco de qual a sua principal utilidade. É comum a gente pensar no direito como talvez uma forma de poder regular tudo, e não é possível fazer isso. Não bastaria a criação de uma norma para determinar que nunca mais ninguém passaria fome no país ou ninguém ficaria pobre porque isso não aconteceria: ao direito o que é do direito, à política o que é da política e às demais ciências o que é para cada espaço. A própria tecnologia tem grande espaço nesse desafio.

E aqui talvez como um ponto teórico, mas sem deixar de ser prático, eu gostaria de trazer um pouco das lições de Lawrence Lessig, um autor que trabalha o tema da tecnologia e da inovação. Ele o faz de modo surpreendente e me parece que as lições dele são muito atuais. Em especial no livro chamado *o Direito ao Código* ou *Código - Versão 2.0*, Lawrence Lessig traz uma ideia, talvez uma alegoria, de que as questões tecnológicas e de profunda inovação, como é o caso da inteligência artificial, deveriam talvez ter um olhar múltiplo. Talvez nós poderíamos olhar para esse objeto a partir de quatro "grandes" grandezas - ou quatro grandezas, para ser menos redundante.

Uma delas, a própria arquitetura tecnológica, ou seja, a tecnologia tem o seu espaço para poder definir os contornos, que impactam inclusive a regulação.

Mas não é só de tecnologia que a gente poderia contar. A tecnologia, de certa maneira, encontra oposição ou certo equilíbrio, cuja outra ponta é seguida ou fortalecida pela própria regulação. Ali, sim, o direito, talvez na sua forma primária, na norma especificamente.

Numa outra ponta, estariam os interesses da sociedade em si. Em contraponto a ela, o do próprio mercado.

Talvez uma figura e uma alegoria que ele traz é uma espécie daquela brincadeira de cabo de guerra, só que com quatro pontas, em que, numa das pontas, estaria a sociedade; na outra ponta, oposta, o mercado; na ponta horizontal, aqui, a arquitetura tecnológica; e na outra, a regulação. E talvez o mundo ideal seria que essas quatro grandezas tivessem potência equivalente, para que a gente mantivesse em equilíbrio, e que a norma fosse um desses componentes, e não o único, para que a gente talvez não colocasse expectativa demasiada na norma, a ponto de achar que a norma resolveria a tecnologia, o mercado e a sociedade.

Essas grandezas, embora independentes, se comunicam, e é essa comunicação é que cria talvez um marco regulatório eficiente, que respeita a própria tecnologia, a arquitetura tecnológica e a sua evolução, mas também respeita o mercado, sem deixar de respeitar e prezando pela sociedade, em especial pelos seus direitos fundamentais.

É esse espaço que eu acredito que a regulação deveria ocupar. Não a totalidade ou a panaceia de todos os desafios tecnológicos, mas parte dela.

Então, a partir talvez dessa premissa teórica, é que eu gostaria de trazer nestes meus próximos minutos, algumas das implicações no campo das eleições, olhando para a inteligência artificial, e também no da desinformação.

Recentemente um dos professores que tem estudado muito o tema e que tem colaborado muito com o país, a respeito desse tema de inteligência artificial, é um colega também do IDP, da Professora, que é o André Gualtieri. Ele recentemente chamou a atenção para as eleições na Coreia do Sul. Então, fazendo jus à minha fonte ali, o que eu observei, ele trouxe a respeito do uso da inteligência artificial nas eleições da Coreia do Sul, com a utilização de avatares para interagir com os eleitores. E, dali, diversos desafios.

Partindo um pouco desse exemplo, eu acho que a gente poderia perceber uma dimensão positiva e democrática da inteligência artificial nas eleições, como a facilidade de comunicação, interatividade e disseminação daquele conteúdo programático, daquela propaganda daqueles determinados candidatos.

Por outro lado, se aqueles que interagem não souberem que se trata de uma inteligência artificial, talvez toda essa dimensão positiva se transforme em negativa, com um engano, com uma forma de manipulação ou até de captação indevida, uma falsidade em si.

Então me parece que existem inúmeras possibilidades em que a inteligência artificial possa facilitar o ideal democrático, de que as mensagens cheguem às pessoas que gostariam de receber aquelas mensagens. Se essas mensagens forem verdadeiras e compatíveis com aquela candidata ou aquele candidato, talvez a gente cumpriu um ideal democrático de que as pessoas que gostariam de eleger pessoas que pensam ou sugerem determinados pontos saibam quem são essas pessoas e qual o número delas antes de irem para a urna depositarem o seu voto.

Por outro lado, se esse conteúdo for criado de modo mais artificial do que a inteligência e chegar como um conteúdo enganoso nesse espaço, a gente tem a potencialidade de um dos fenômenos que até o momento a gente não conseguiu enfrentar ou pelo menos resolver. Talvez em nenhum lugar do mundo o fenômeno da desinformação tenha sido realmente resolvido, e um dos motivos, eu acredito, é em razão da multilateralidade desse fenômeno.

Talvez nós não conseguiremos resolver esse fenômeno apenas com regulação ou apenas com punição, e seria necessário um conjunto de mais medidas, desde o espaço de educação, de prevenção até, também, o de repressão, porque, afinal, não dá para brincar com a democracia e esperar que todo mundo se eduque para que a gente possa ter um ambiente digital mais saudável, em especial do ponto de vista democrático.

Olhando especificamente para a aplicação da inteligência artificial na desinformação, gostaria de chamar a atenção para quatro pontos. A inteligência artificial...

*(Soa a campanha.)*

**O SR. DIOGO RAIS** - Tenho mais um minutinho? Obrigado.

A inteligência...

Na verdade, há 45 segundos que eu estou preocupado com a voz de que o Senador falou. *(Risos.)*

Então, ela impacta na quantidade de desinformação, ou seja, a inteligência artificial não descansa, não dorme e pode produzir aquilo numa quantidade demasiada, tanto para o bem quanto para o mal.

Além disso, a própria qualidade do teor. O conteúdo pode ser muito mais sofisticado. Imagine um vídeo no qual você apareça falando ou fazendo algo que nunca fez. Imagine como é difícil dizer para alguém: "Você acredita em mim ou no que os seus olhos veem?". Veja como é injusto! E, de uma certa forma, essa sofisticação também é potencializada pela inteligência artificial.

E os últimos dois pontos que eu gostaria de mencionar também são a forma, sobre você ter um roteiro cada vez mais eficiente, considerando os dados disponíveis, e, por fim...

*(Soa a campanha.)*

**O SR. DIOGO RAIS** - ... o direcionamento desse conteúdo.

A inteligência artificial pode criar talvez um dos maiores temores: a desinformação ou a *fake news* sob medida, aquela mensagem adequada a determinado grupo de pessoas, podendo promover um gigante impacto.

Gostaria de agradecer mais e pedir desculpas por ter avançado um pouquinho no tempo.

Mas continuo à disposição de todos.

Muito obrigado, Senador.

Obrigado a todos.

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado, Sr. Diogo Rais, Diretor do Instituto Liberdade Digital. Parabéns pela exposição!

Eu anotei aqui algumas das partes da fala para comentários depois, mas realmente é uma tarefa muito difícil - imagino. Primeiro, porque a natureza, quando se fala de eleição, a natureza das campanhas, e eu posso dizer isso por ter sido eleito através de uma campanha, mas existem coisas que são boas e coisas ruins durante a campanha, descontando-se a tecnologia.

O que eu imagino que a população quer? O que um eleitor quer? Ele quer é conhecer quem é o candidato, quer saber da vida do candidato, quer saber das ideias, dos valores, das propostas. São as coisas positivas, mas, infelizmente, a gente vê também, descontando, sem a tecnologia, a gente vê muito ataque, muitas coisas, muitas mentiras, muitas vezes, entre os próprios candidatos, o que é muito ruim. Eu sempre detestei isso, muito antes de ser político, e a máquina pode, a inteligência artificial pode multiplicar esses efeitos, sem dúvida nenhuma.

E uma dificuldade que sempre existe - eu estava imaginando quando você estava falando das quatro cordas ali - é o que separar, por exemplo, em termos de conteúdo: o que é *fake news*, realmente; o que é opinião, porque a pessoa pode acreditar, ela pode até acreditar em uma coisa errada, mas pode acreditar, e pela liberdade ela tem o direito de expressar a sua opinião. Agora, é difícil fazer essa separação.

Outra coisa importante são os novos conhecimentos. Deixe-me explicar isso de uma forma um pouco... Vamos usar a história. Vamos pensar assim: o Giordano Bruno, lá atrás, foi queimado na fogueira porque falava que a Terra não era o centro do universo. O conhecimento da época, o que era acreditado como conhecimento - e que ninguém podia ser contra - era que a Terra era o centro do universo. Como ele veio com uma ideia nova, na época talvez uma *fake news*, não era bem uma *fake news*, mas era um novo conhecimento, aquilo, então, foi interpretado dessa forma. Ou seja, a gente tem que tomar muito cuidado.

Ou mesmo Einstein, que foi pensar, lá atrás, quando ele propõe uma teoria geral ou restrita da relatividade, contrária... não contrária, mas expandindo o conhecimento de Newton - o *Principia*, de Newton. De repente, ele era um cara apresentando uma coisa completamente nova que, entre aspas, "desmentia"...desmentia, não, colocava dúvidas no que já estava estabelecido.

A gente tem que ter a liberdade também de apresentar essas coisas. E, ao mesmo tempo, também outro ponto importante são as ofensas - não é? -, exatamente as ofensas, que são terríveis - os *haters*, que a gente vê na internet.

Eu tenho discutido bastante isso aqui também. Eu acredito muito na necessidade de autoria. Eu quero dizer assim: não importa se você... Eu tenho o meu CPF, eu sou homem, mas eu posso me chamar ali dentro da internet como um nome feminino, Rosa de alguma coisa. Mas o meu CPF tem que estar registrado com aquilo, de forma que a gente saiba exatamente, a todo o tempo, quem está por trás daquilo que está sendo colocado. A liberdade implica também responsabilidade sobre o que você... Bom, são umas partes aí que eu...

Mas parabéns pela apresentação!

Na sequência, eu gostaria de passar a palavra para a Sra. Celina Bottino, que está remota, Diretora de Projetos do Instituto de Tecnologia e Sociedade.

Sra. Celina, a senhora tem dez minutos. Por favor, controle o tempo aí pelo seu lado.

Obrigado.

**A SRA. CELINA BOTTINO** (Para expor. *Por videoconferência.*) - Muito obrigada.

Eu vou tentar aqui passar aqui uns eslaides, vamos ver se vai funcionar. *(Pausa.)*

Estão vendo?

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Sim. Está aqui na tela.

**A SRA. CELINA BOTTINO** *(Por videoconferência.)* - Bom, primeiro bom dia a todos.

Meu nome é Celina Bottino, sou Diretora de Projetos do ITS, que é o Instituto de Tecnologia e Sociedade. É um prazer estar aqui representando o ITS, porque este ano completamos dez anos nesta missão de contribuir com um debate informado sobre os impactos da tecnologia dentro da sociedade. Eu tenho também acompanhado as audiências e quero parabenizar aqui a Comissão Temporária pela organização de todas essas conversas, que são uma oportunidade de a gente incluir

outras vozes e aprender com pessoas diferentes, históricos e atuações sobre os impactos da inteligência artificial e também de como a gente consegue avançar nesse desafio regulatório.

Bom, primeiro, complementando, espero ou imagino que as falas serão diferentes, mas eu vou focar a minha contribuição no tema das *deepfakes* e seus potenciais impactos aí nesse contexto eleitoral.

Mas eu queria começar com uma frase aqui, uma citação do Ian Goodfellow, que basicamente diz que no futuro as pessoas não deveriam acreditar mais nas imagens e nos vídeos que eles veem na internet. E o Ian foi quem criou o algoritmo por trás das *deepfakes*. Talvez ele esteja certo.

Vamos ver aqui como é que a gente... A ideia desta conversa é tentar explicar um pouquinho, por isso é que eu vou dividir a minha fala em três momentos. Primeiro, explicar sobre a evolução e a criação e do uso de *deepfakes*. Depois, rapidamente passar um pouquinho sobre o uso dessas *deepfakes* no contexto de guerra e eleição, contextos envolvendo aí questões democráticas. E, por fim, passando pelas iniciativas que têm sido consideradas para lidar com esse desafio.

Então, para começar, eu queria só trazer um pouco do histórico, de onde vem a tecnologia por trás dessas *deepfakes*. Como eu comentei, foi criada aqui pelo Ian Goodfellow, com alguns colegas, que basicamente é grátis.

Primeiro, vou dar um passo atrás. *Deepfake* seria basicamente manipular e criar imagens para que elas fossem muito parecidas com a realidade. Tentar manipular algo para se tornar real, para parecer real. E isso é possível graças a essa *generative adversarial network*. O que é isso? São basicamente duas redes neurais, de duas IAs, em que uma cria uma imagem, com base em um banco de dados, ela gera uma imagem, e a outra é usada para identificar as falhas naquela imagem. Uma vez que identifica, ela manda essa mensagem, daí ela automaticamente tenta melhorar onde precisa ser melhorada.

Então, na verdade, isso foi criado entre os amigos conversando, que estavam tomando uma cerveja lá num bar no Canadá, e esses amigos do Ian comentaram, "a gente está tentando criar uma possibilidade de a IA criar imagens sozinha". Só que estava ruim, não estavam conseguindo. Aí ele pensou "se a gente botasse os dois juntos?" Enfim, daí foi isso que gerou. Isso então foi em 2014, então não tem tanto tempo assim.

E, no início, a gente vê que realmente ainda era muito rudimentar, claramente, facilmente identificado como algo que não era real, mas a tecnologia foi evoluindo. Em 2016, já era possível manipular também um pouco a imagem, colocando óculos, por exemplo, e alguns outros acessórios, mas mesmo assim se via ainda claramente que não eram imagens realistas e verdadeiras.

O salto de qualidade aconteceu em 2017, quando a Nvidia identificou, aplicando esse algoritmo, começou a ver que se ela fizesse de forma gradativa, ou seja, começasse uma imagem com uma resolução mais baixa e gradativamente fosse melhorando, a imagem ia ficando melhor. Então esse aqui é o resultado de pessoas, fotos artificialmente criadas. As pessoas não existem, mas foram pessoas criadas para parecerem pessoas famosas, com base em bases de dados de pessoas, enfim, de artistas famosos.

Então a gente já vê realmente uma grande diferença. Em dois anos, você vê como já deu um salto de qualidade. E só lembrando que isso não significa só em questão de imagem. A tecnologia pode ser usada para fazer objetos, como aqui carros, ambientes, enfim.

E a tecnologia foi sendo evoluída, foi colocando possibilidade de dublagem e tudo mais, até que ela tomou mais os holofotes quando, em 2019, começou a ser usada de forma mais substancial em *sites* de pornografia. Então as imagens de artistas, e o caso aqui da Scarlett Johansson, que teve a sua imagem trocada, num vídeo de pornografia, e daí isso ganhou os holofotes e discussões sobre os potenciais usos dessa tecnologia.

A gente sabe que os problemas e efeitos podem ser muitos: de ordem psicológica, pode-se criar situação de intimidação, de *bullying*; na ordem financeira, podem-se fazer fraudes; e, também, há impacto na sociedade como um todo, quando se mexe na informação em contexto de eleição e em guerras, por exemplo, que acho que é um tema sobre o qual a gente vai se debruçar.

Esse é um assunto que começou a ser discutido especialmente nos Estados Unidos, desde 2018. Foi quando esse comediante, o Jordan Peele, criou esse vídeo no qual o Obama insultava o então candidato Donald Trump. Na época, gerou-se uma certa confusãozinha, mas ficou claro depois que era o trabalho desse comediante - mas isso já acendeu o alerta do possível uso dessa tecnologia em um contexto eleitoral.

Passando aqui mais para o... Enfim, desde então, nas eleições de lá de 2020, os Estados Unidos se debruçaram sobre esse tema e daqui a pouquinho a gente vai ver um pouquinho quais foram as iniciativas em relação a essa discussão lá nos Estados Unidos.

Mais recentemente, no contexto da Guerra Ucrânia-Rússia, surgiu esse vídeo do Presidente Zelensky, que, a princípio, pedia para o povo entregar as armas, e que foi claramente manipulado, mas foi de forma muito rápida identificado como algo falso.

Enfim, isso são dois resumos muito pontuais, só para dar uma ideia do potencial efeito danoso, a depender de como for usado. É claro que no Brasil e em outros lugares tem outros exemplos, mas acho que agora seria mais interessante a gente já passar para o que fazer. A ideia agora é fazer um panorama geral aqui de algumas iniciativas pelo mundo afora.

Então, começando pelas iniciativas legislativas: nos Estados Unidos, como eu comentei, dois estados já aprovaram leis proibindo o uso de *deepfakes* no contexto eleitoral - seriam a Califórnia e o Texas - e, no âmbito federal, tem-se discutido também um projeto de lei para criminalizar a produção de *deepfakes* que não tenham uma marca d'água identificando-o como tal.

A China também lançou uma ofensiva grande contra o uso de *deepfakes*, botando medidas dizendo que esses conteúdos têm que ser rotulados como tal, e, além do mais, teria que ter o consentimento da pessoa antes de ter sua voz ou imagem alterada. Além disso, o Governo criou uma lista de 41 empresas que estariam de acordo com os requisitos da legislação.

O Canadá vem atuando nessa frente por três vieses: primeiro, prevenção, estão fazendo campanhas grandes de conscientização; detecção, estão investindo em tecnologia para aperfeiçoar formas de identificar; e, por fim, a resposta, que seria também uma legislação no sentido de criminalizar a produção de *deepfakes* que tenham essa intenção de causar confusão ou dano.

Na União Europeia, a gente tem medidas no DSA - no *Digital Services Act* - também para lidar com isso; e, no AI Act, também tem, no art. 52, 3, uma obrigação de informar como conteúdo foi artificialmente gerado.

Na Coreia do Sul, também tem uma lei tornando ilegal a distribuição de *deepfakes* com esse intuito de causar dano.

Então, a gente sabe que todas essas leis têm essa questão de causar dano ao interesse público como um elemento importante para caracterizar a sua criminalização.

No Brasil, já tem o PL também, o 1.002, que criminaliza, muda o Código Eleitoral para criminalizar o uso de *deepfakes* em período eleitoral; e, no PL 2.338, também tem obrigações de transparência. Então, todas vão nesse sentido de trazer mais informação para o público.

Outras iniciativas, além das legislativas, seriam, por exemplo, a de Taiwan, que vive sendo um bombardeado com notícias falsas no contexto da relação deles com a China. Lá parece que estão criando uma "imunidade de *nerd*". Esse é o termo para treinar as pessoas e "checadores" para identificar esse conteúdo que foi gerado artificialmente. O Japão também teve a iniciativa de criar um perfil originador...

(Soa a campanha.)

**A SRA. CELINA BOTTINO** (Por videoconferência.) - ... para dar a originalidade do conteúdo.

Já estou terminando.

Essas outras iniciativas também têm a ideia de mostrar de onde vem o conteúdo e há outras formas de tecnologia para detectar quando você está falando de *deepfake*, mas, enfim, nem tudo é negativo, então, esses são usos positivos, ajudando pessoas a voltar a poder falar usando comando de voz, a poder contar a história do Holocausto...

A gente não pode dizer que tudo é negativo, mas, resumindo, não teria uma solução única, não tem uma bala de prata nem todo uso é negativo. Talvez, a lei só, como o Diogo falou, fosse ineficiente, teria a sua limitação. Talvez o Brasil já tenha alguns elementos, aqui na nossa legislação, que dê conta disso, mas voltando, acho que o ponto mais importante seria campanhas de educação midiática, para ensinar o olhar crítico sobre esse conteúdo.

Então, termino aqui. Muito obrigada pela atenção.

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado, Sra. Celina Bottino. Parabéns pela apresentação bastante esclarecedora e bastante interessante. Eu fiz uma anotação aqui, ao longo da sua fala, primeiro, sobre a questão da prevenção como educação, também extremamente importante, e sobre as leis, proibindo também, mas, logicamente, sempre vai sobrar um percentual daquelas pessoas ou instituições que fazem porque querem fazer errado realmente, não é?

Aí, para jogar uma ideia para se pensar, da mesma forma que nós temos vírus na internet, *malwares* e uma série de coisas que têm que ser combatidas o tempo todo por *cybersecurity*, segurança cibernética, é uma batalha constante para se proteger dados, proteger a integridade e a confidencialidade, proteger o funcionamento dos sistemas, eu imagino que, em um futuro não muito distante, possivelmente já até os próprios sistemas de inteligência artificial vão ter que ser criados com defesas para que identifiquem, rapidamente, o que é verdade e o que não é, em termos de imagens falsas, vozes falsas,

vídeos falsos, como se fosse uma segurança cibernética contra isso aí na inteligência artificial. Não sei se já está sendo criado, mas é, possivelmente, um dos caminhos que vão se seguir para segurar essa situação.

Na sequência, eu concedo a palavra à Luciana Moherdaui, Pesquisadora do Grupo Jornalismo, Direito e Liberdade, da Escola de Comunicação e Artes da Universidade de São Paulo. Professora, a senhora tem dez minutos. Por favor.

**A SRA. LUCIANA MOHERDAUI** (Para expor. *Por videoconferência.*) - Bom dia a todos.

Agradeço pelo convite e também quero cumprimentar a Comissão pela iniciativa das audiências públicas. Este é um tema bastante importante que não se esgota assim tão facilmente.

Eu vou falar sobre a iniciativa de ontem do Presidente Biden a respeito da inteligência artificial e sobre como os jornais lá têm olhado para essa questão - a imprensa americana tem considerado as eleições de 2024 as primeiras orientadas por inteligência artificial -; e, na sequência, se eu não estourar o tempo, vou fazer uma pequena abordagem sobre o impacto da IA no jornalismo.

Então, como é sabido - a imprensa brasileira também noticiou -, o Presidente Biden anunciou ontem uma ordem executiva que impõe novas obrigações de segurança aos desenvolvedores de inteligência artificial e apela a uma série de agências federais para mitigar os riscos da tecnologia. Essa ordem exige que as empresas que desenvolvem sistemas de IA mais avançados realizem testes de segurança e notifiquem o governo dos resultados antes de lançarem os seus produtos. A ordem tem recomendações - não é uma exigência, é uma recomendação - para que fotos, vídeos e áudios desenvolvidos por tais sistemas tenham marcas d'água, para deixar claro que foram criados por inteligência artificial. Segundo o *The Washington Post*, este é um temor nos Estados Unidos, em relação à eleição do ano que vem: que essas ferramentas... algumas já estão na rua, mas que elas desestabilizem a campanha eleitoral. A ordem também insta o Congresso a aprovar legislação bipartidária sobre privacidade; é um tema que está sendo discutido lá, a privacidade para proteger os americanos, especialmente as crianças. Como sabemos, não existe uma lei federal nos Estados Unidos de privacidade.

A ordem do Biden vem no momento em que haverá um grande debate sobre inteligência artificial no Reino Unido. O *The Washington Post* conta que, para essa ordem executiva, o Presidente Biden olhou para a lei europeia, olhou para a China, para Israel, enfim, para montar esse documento e apresentar.

Desde que a popularidade do ChatGPT explodiu no ano passado, legisladores e reguladores têm... todo mundo tem debatido sobre como a inteligência artificial - e o Senador Pontes já falou aqui no início, na abertura - pode alterar empregos, espalhar desinformação e potencialmente desenvolver seu próprio tipo de inteligência. As novas regras, segundo o *The New York Times*, que já começam a valer, devem entrar em vigor nos próximos 90 dias, devem enfrentar muitos desafios, alguns legais, outros políticos; mas, enfim, a ordem visa a esses sistemas mais avançados em grande parte. E, no caso ali de marcas d'água, eu penso que é um avanço, mas também é um problema, porque as campanhas eleitorais não se dão só no âmbito oficial de uma campanha, então a militância, os apoiadores, os influenciadores ali precisavam encontrar uma maneira de isso ser... Claro que isso é impossível de se excluir completamente, mas de dirimir aí esse tipo de conteúdo. Também é verdade que muitos candidatos... Inclusive a imprensa tem feito muitas reportagens sobre isso, o *The Washington Post* principalmente, mostrando que campanhas pequenas com pouco financiamento podem usar sistemas de IA para arrecadação, envio de *e-mails* aos seus eleitores, enfim, mas, por outro lado, também pode haver pedido de arrecadação enviado por sistemas que são pensados para burlar as regras.

Um colunista do *The Washington Post* recentemente fez uma alerta sobre as eleições do dia 2024, eu acho que é uma coisa que parece uma utopia: ele sugere que os candidatos se comprometam a não usar IA para enganar os eleitores. Então, nos Estados Unidos, sabemos, por exemplo, que o Presidente Trump não vai se comprometer em relação a isso. Então, é um apelo que tem sido feito, que também não acho que isso, de novo, não vai ter um efeito somente em campanhas, mas muitos trabalham com apoio, não é.

Uma pesquisa que o *The Washington Post* divulgou recentemente mostra que 85% dos americanos disseram estar muito ou um pouco preocupados com a propagação de vídeos ou áudios enganosos produzidos por meio de inteligência artificial e 78% estavam preocupados com o fato de a inteligência artificial contribuir para difusão de propaganda política enganosa. Então, por exemplo, a campanha do Governador da Flórida, DeSantis, compartilhou uma imagem falsa gerada por IA do ex-Presidente Trump abraçando o Anthony Fauci. Esse abraço nunca aconteceu. Nas primárias para Prefeito em Chicago, conta o *The Washington Post*, alguém usou inteligência artificial para clonar a voz de um candidato chamado Paul Vallas em uma reportagem falsa, fazendo parecer que ele aprovava a brutalidade policial. Ambos os casos poderiam ser desmascarados, mas o que acontece - pergunta o colunista - quando uma imagem ou clipe de algo chocante se tornar viral em um estado decisivo, por exemplo, pouco antes de uma eleição, que tipo de caos ocorrerá quando alguém usar um *bot* para enviar mentiras personalizadas a milhões de eleitores diferentes?

Eu vou avançar um pouco, senão vai ficar muito longo.

Claro, não há nada de ruim em usar a inteligência artificial, é algo que não pode ser impedido, nenhum avanço tecnológico, mas essa é uma preocupação real nos Estados Unidos para 2024.

Vou fazer o corte para o jornalismo, que também é o tema desta audiência pública, o uso da inteligência artificial para o jornalismo. No estágio atual em que estão esses *chats*, alguns têm até *links* para saber mais, enfim, não creio que seja recomendável para apuração e até mesmo para redação de textos. Eu passei um semestre como voluntária na Faculdade de Arquitetura e Urbanismo da USP acompanhando alunos da disciplina Cultura Urbana.

Eles se revezaram entre aulas teóricas e pelo ChatGPT. O desafio era: a partir das aulas teóricas, produzir artigos. Então os artigos, durante todo o semestre, e também imagens, porque eram alunos de design, os artigos todos padronizados, com informações que não poderiam ser checadas ali. O programa usado foi o ChatGPT, o ChatGPT grátis. Muitas imagens também de inteligência artificial, que foram criadas a partir de contextos históricos também. Apresentaram produtos que não eram...

*(Soa a campanha.)*

**A SRA. LUCIANA MOHERDAUI** *(Por videoconferência.)* - ... não faziam referência ali ao que foi feito em sala de aula. Então é uma responsabilidade muito grande. Foi importante esse método ali na Faculdade de Arquitetura. A Profa. Giselle Beiguelman, que trabalha com isso há muito tempo, é autora de livros, pensa muito a arte a partir da inteligência artificial, ali mostrou que é preciso ter muito cuidado.

A tecnologia tem que avançar, não sou contra, mas assim, em algumas aplicações. Para o jornalismo, ainda não. Há algumas iniciativas, como a da Bloomberg, para usar a inteligência artificial para a sua base de dados. Isso é algo em que eu não vejo problema algum, mas isso não impede a apuração do jornalista, não impede ali o jornalista no *tête-à-tête*, fazendo o seu trabalho de campo, não é? Um grande exemplo disso é...

*(Soa a campanha.)*

**A SRA. LUCIANA MOHERDAUI** *(Por videoconferência.)* - ... o colunista da *Folha*, Ruy Castro. Ele é um dos críticos do uso para apuração, até porque sabemos como é o trabalho dele, como é que ele faz a apuração.

Então, com isso, eu encerro a minha participação. Mais uma vez, agradeço o convite, Senador, e agradeço à Comissão pela iniciativa.

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado. Obrigado, Profa. Luciana Moherdaui, pesquisadora do grupo Jornalismo, Direito e Liberdade da Escola de Comunicação e Artes da Universidade de São Paulo. Muito obrigado pela sua apresentação. Parabéns.

Eu passo agora a palavra à Sra. Patricia Blanco, Presidente do Instituto Palavra Aberta. A senhora tem uns dez minutos só.

**A SRA. PATRICIA BLANCO** (Para expor.) - Obrigada, Senador. É um prazer enorme estar aqui e uma honra poder contribuir com esta Comissão Temporária de Inteligência Artificial.

Eu tenho discutido muito esse tema, inclusive no âmbito do Conselho de Comunicação Social aqui do Congresso, no qual eu tenho a honra de estar Vice-Presidente neste mandato. Então, para mim é realmente muito importante a gente discutir ativamente e dentro dessas audiências públicas, que são fóruns qualificados de participação da sociedade civil e de todos os agentes envolvidos nesse tema, que é tão complexo, não é? Como o Prof. Diogo colocou, é um tema do nosso tempo, não é? Estamos aqui vivenciando isso, essa mudança tecnológica.

Eu queria também agradecer o convite do Senador Eduardo Gomes e do senhor, Senador Astronauta Marcos Pontes.

Queria saudar a presença da Samara Castro aqui, da Secretaria de Políticas Digitais, que, junto com o Secretário João Brant, que está aqui também acompanhando e vai participar, estão fazendo um trabalho brilhante nessa interlocução da sociedade civil com o Governo, principalmente no eixo da educação midiática e da defesa da liberdade de expressão.

Queria saudar os meus parceiros aqui nesta audiência pública, que fazem com que a gente tenha realmente um debate extremamente qualificado. O Prof. Diogo, meu companheiro aqui de lutas no combate à desinformação, a Profa. Tainá; a Renata Mielli, que está agora como coordenadora do CGI; Marcelo Rech, Presidente da ANJ; Celina Bottino, do ITS, a Profa. Dora Kaufman, que nos agradeceu com a participação num livro lançado pelo Palavra Aberta recentemente, e a Profa. Luciana, que trouxe um ponto fundamental sobre essa questão do artigo do *Washington Post*. Profa. Luciana, você roubou o meu artigo, eu ia citá-lo, viu? Mas eu vou citá-lo mesmo assim.

Bom, queria também fazer alguns agradecimentos ao Prof. Gustavo Borges, do LabSul, de Santa Catarina, que fez alguns apontamentos para a minha fala de hoje; à Mariana Ochs, Coordenadora de Educação Midiática, uma das coordenadoras do Instituto Palavra Aberta, que tem feito um estudo exaustivo sobre o impacto da inteligência artificial em questões relacionadas à educação e sobre a importância de refletirmos, a partir do eixo da educação midiática, para combater e mitigar os danos causados pelo uso inadequado da inteligência artificial, e ao Alexandre Sayad, que foi o nosso parceiro para a divulgação e publicação deste livro, Senador, que eu gostaria de deixar com o senhor, *Inteligência Artificial e Pensamento Crítico: Caminhos para a Educação Midiática*, que foi uma contribuição feita a partir da tese de mestrado do Prof. Alexandre e que tem o prefácio da Profa. Dora Kaufman.

O ponto fundamental nessa questão que o Alexandre traz nesse debate é o pensamento crítico. O que nos torna humanos nessa discussão em relação à inteligência artificial? Quais são essas habilidades que nós temos que desenvolver na medida em que a tecnologia acaba se sobressaindo, acaba impactando todas as nossas relações, sejam relações sociais, sejam relações humanas, sejam relações do nosso dia a dia? Hoje a tecnologia já está presente, a inteligência artificial já está presente em todas as nossas relações. A pergunta que a gente tem que considerar é: como nós queremos construir, qual é o tipo de sociedade que queremos constituir a partir do uso incessante, intenso e cada vez mais presente da tecnologia nas nossas vidas? Qual é o tamanho do espaço para o humano? E a gente tem que pensar numa legislação ou até no momento em que a gente discute se a inteligência artificial vai nos substituir ou não.

Um ponto, dentro dessa questão do que nos faz humanos, é justamente pensar nos princípios que podem e que devem nortear qualquer legislação que trate sobre inteligência artificial.

Analisando o projeto do Senador Rodrigo Pacheco, que já está em discussão, eu queria trazer alguns princípios que já constam do art. 3º do PL, que são a questão da autodeterminação, da liberdade de expressão, dos direitos humanos, da criatividade, da interpretação, da transparência. Eu acho que a Celina colocou bem essa questão da importância e da necessidade da transparência no uso da inteligência artificial.

Um ponto que eu queria incluir nesse artigo - deixo como sugestão de inclusão - é justamente a questão da análise crítica do uso dessa ferramenta e de como a gente pode melhorar essa utilização da tecnologia a partir do pilar transparência, colocando, claramente, identificando claramente, que determinados assuntos ou determinadas funções estão sendo mediados por uma tecnologia, com o aprendizado de máquina, com inteligência artificial, com depuração de dados e seguindo também outras questões que, para mim, são fundamentais, que tratam dos direitos de todos nós, dos cidadãos, a partir dos pontos trazidos no texto.

Há alguns pontos que eu queria trazer aqui como fundamentais também.

No Capítulo II, art. 5º, trata do direito do cidadão: informação prévia, explicação, contextualização, determinação, participação humana, privacidade e - um ponto que para mim é fundamental - não discriminação. E, nesse ponto específico da não discriminação, a gente tem que tentar olhar e analisar de que forma essas ferramentas de inteligência artificial estão sendo criadas a partir de vieses já discriminatórios da nossa sociedade e que podem impactar grupos minoritários e até grupos na sociedade e reforçar direitos e impactar questões de vieses e de fundamentos até abusivos em relação a isso.

Então, a gente precisa pensar na arquitetura e no *design* dessas ferramentas, levando em consideração, de novo, a questão do *do no harm*, quer dizer, não causar dano. Como essa tecnologia precisa ser pensada, precisa trazer esses pontos de que toda tecnologia vem, pode causar um bem enorme, mas também pode causar dano.

Então, como é que a gente pode utilizar justamente a legislação para mitigar os danos possíveis que podem ser causados a partir dessas tecnologias?

E dois pontos fundamentais que eu vejo que estão sendo pouco explorados no PL. O primeiro é no que tange à educação. Embora já conste do art. 2º, nos fundamentos, acesso à informação e à educação, não há nenhuma outra menção a esse termo no PL, quer dizer, na proposta toda.

Então, temos que ter na educação, principalmente na ótica da educação digital e midiática, a base para tornar o cidadão mais resiliente aos efeitos adversos que a IA e a ação algorítmica podem causar na nossa sociedade.

Então, precisamos trabalhar, sim, esse ponto que é fundamental, e tendo a educação que precisa promover um termo que a Profa. Letícia Cesarino, que está hoje no Ministério dos Direitos Humanos, cravou, muito importante, que é a desalienação técnica, que é justamente a utilização da IA do ambiente humano mesmo, trazer o humano para sair fora dessa alienação algorítmica, desse impulsionamento de dados tão focados e tão estratificados a ponto de você poder fazer uma comunicação, e aí, no caso das eleições, muito específica, com micro *target* naquele público, quase que uma personalização da informação que chega até o eleitor.

E o segundo ponto que também eu gostaria de ver um pouco mais claro na legislação é o da IA como ferramenta de contenção de danos, para proteger justamente essa utilização dessa tecnologia tão fantástica que chega até nós. E quero lembrar que essa mesma tecnologia que pode causar danos deve ser usada para mitigá-los. Como exemplos: criação de marcas d'água; imagens geradas por inteligência artificial; indicação do uso da tecnologia; reforço do uso seguro; preservação da integridade da informação...

*(Soa a campanha.)*

**A SRA. PATRICIA BLANCO** - ... a partir da análise de padrões; atitude proativa na detecção de conteúdos adulterados; validação, verificação, autenticidade, que podem ser incorporadas às ferramentas de IA, promovendo o uso de forma preventiva; soluções de ferramentas de IA que podem e devem - eu gosto de reforçar bem isso -, devem ser usadas como uma proteção, um escudo mesmo de proteção para conteúdos fraudulentos.

E, só para encerrar, Senador, eu queria trazer uma questão que não está na legislação, mas é uma questão de cunho filosófico até, um pouco, que é a questão ética. Que é: como que a gente trata o uso ético de qualquer ferramenta ou de qualquer tecnologia? E aí vem muito essa questão da necessidade dos pactos coletivos em relação ao uso das inteligências artificiais e dessas ferramentas. Como que a gente garante que o impacto de decisões automatizadas favoreça os direitos individuais e o bem comum?

*(Soa a campanha.)*

**A SRA. PATRICIA BLANCO** - Como que a gente ressignifica e avalia se questões como liberdade, direitos humanos, justiça e equidade estão presentes e sendo valorizados? A tecnologia deve ser construída sobre valores de transparência, justiça, equidade, privacidade, segurança e confiabilidade.

E um ponto que eu sempre gosto de trazer é que todo o uso da tecnologia ou de qualquer ponto em relação a ferramentas que estão disponíveis, seja para *deepfake*, seja para proliferação de desinformação, esbarra num ponto que é a responsabilidade de todos nós enquanto cidadãos, para que isso torne realidade. E aí é aquela questão: a gente não vai conseguir fazer pactos coletivos ou pactos de não agressão e de não utilização de ferramentas de tecnologia, de *deepfakes*, para combater a desinformação no período eleitoral? Estamos preparados para isso?

*(Soa a campanha.)*

**A SRA. PATRICIA BLANCO** - Essas são as minhas contribuições. E, de novo, quero trazer a questão do olhar humano e cidadão para qualquer uso da tecnologia.

Muito obrigada.

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado, obrigado. Parabéns pela apresentação, Sra. Patricia Blanco, Presidente do Instituto Palavra Aberta. Aliás, a senhora trouxe um tema interessante sobre o ser humano. Eu tenho falado aqui várias vezes: eu trabalho com tecnologia a minha vida inteira, mas uma coisa importante sempre a ter em conta é que o ser humano tem que ficar no centro, senão não faz sentido. Então, esse é um ponto importante.

E eu acho que essa discussão que a gente tem tido sobre inteligência artificial, os aspectos, os riscos principalmente, tem levantado também uma necessidade muito importante - a senhora falou na questão de ensinar as pessoas sobre isso -: a discussão da própria natureza humana, que não é de hoje, e descontando a inteligência artificial, a natureza humana de mentir, enganar. Obviamente isso já existe, isso faz parte da natureza e isso pode e deve ser melhorado aos poucos, através da cultura, educação e melhorando essa convivência humana, tanto no aspecto intrapessoal quanto interpessoal também. E isso acaba refletindo na inteligência artificial, nas nossas preocupações do uso da inteligência artificial, que certamente é a tecnologia criada por humanos, os bancos de dados são criados por humanos, buscam dados que vão ser utilizados para a inteligência artificial e ela amplia essas qualidades ou defeitos humanos - isso é importante.

Quando a gente fala de discriminação, por exemplo, que pode ser tanto no desenho, mas principalmente no uso do banco de dados, às vezes um banco de dados parcial, ou parcial nos dois sentidos de parcial, parcial de ser um pedaço só, parcial no sentido de ter algum viés, isso aí vai ter como resultado um sistema que aprendeu errado. Então, é importante sempre ter isso em mente.

Outra coisa importante sobre as decisões, outra coisa que eu tenho ressaltado aqui é que a inteligência artificial certamente tem uma capacidade muito grande de análise de dados, de resumo de dados, de classificação e aplicação específica, mas eu não gosto da ideia. Lógico que existem sistemas autônomos que a gente tem que discutir depois, como carros, aviões, etc., mas se está falando de um certo tipo de decisão. Eu acho que as decisões sobre as vidas das pessoas têm que ser

tomadas pelas pessoas, não por uma inteligência artificial, mesmo porque, se a gente for pensar, nem sempre o lógico é justo. Só para a gente ter em mente isso aí.

Bom, vamos prosseguir.

Eu gostaria de chamar agora, para a sua apresentação de dez minutos, o Sr. Marcelo Antônio Rech, Presidente da Associação Nacional de Jornais.

Sr. Marcelo, concedo-lhe a palavra por dez minutos e peço-lhe para controlar o tempo por lá.

Obrigado.

**O SR. MARCELO ANTÔNIO RECH** (Para expor. *Por videoconferência.*) - Estou controlando aqui.

Bom dia, Senador. Bom dia a todos os colegas.

Parabéns pelas manifestações anteriores, são muito interessantes.

Eu gostaria de dividir minha fala aqui rapidamente em três grandes pontos.

Aqui eu falo também, além de Presidente da Associação Nacional de Jornais, como membro do Comitê Executivo da Associação Mundial de News Publishers, que é uma entidade que reúne milhares de jornais e veículos de comunicação em todo o mundo.

O tema da inteligência artificial é um tema central na comunicação e no jornalismo já há um bom tempo e, neste ano, seguramente, se tornou o principal eixo de discussão em todos os fóruns internacionais sobre as consequências e as implicações da IA.

Um dos primeiros pontos importantes que eu gostaria de ressaltar é que nós estamos fazendo um esforço muito grande para chamar a atenção de todos os *players* mundiais para que a gente não cometa os mesmos erros que cometemos nos primórdios da era digital da internet, em 1994 e em 1995, quando o deslumbramento e o encantamento, justificados - e, de certa forma, até como nós estamos vivendo agora nessa nova era da inteligência artificial -, cegaram todas as instituições, digamos, com responsabilidade de estabelecer alguma regulação ou alguma ordem nisso, à luz de um certo anarquismo digital ou de um desenvolvimento quase que anárquico, supostamente no bem para toda a humanidade. Em grande medida, sim, mas, em outra parte, nós vemos os efeitos colaterais desse desenvolvimento se manifestando hoje nas nossas vidas, nas relações sociais e nas relações geopolíticas de todo o planeta. Então, que não se cometam os mesmos erros, tanto do ponto de vista de transparência, de não regulação e, por fim, de concentração.

Então, nós temos visto que quem desenvolve a inteligência artificial hoje são os mesmos *players* que têm hoje os impérios digitais altamente concentrados ou são os novos impérios que estão se construindo, reduzindo não só a possibilidade de uma multiplicação de desenvolvedores, de um modelo central para a economia e para a vida dos cidadãos, mas como também concentrando, perigosamente, em poucas mãos, uma quantidade absurda de poder - eu não falo só do Brasil, eu falo do mundo. Esta é uma preocupação central de todos os governos da Europa e, enfim, do próprio Biden ontem: estabelecer um mínimo de regulamento para que essa concentração não ameace a segurança nacional - por exemplo, a segurança nacional dos Estados Unidos -, mas as seguranças de cada país.

O segundo ponto é, especificamente, sobre direitos autorais no jornalismo. Os desenvolvedores de inteligência artificial não se caracterizam pela transparência. A gente não sabe, ninguém sabe exatamente de onde vieram e de onde vêm essas informações, de onde vêm esses conteúdos, e, propositalmente, eles não informam no detalhe, mas a gente sabe que um modelo como o ChatGPT, por exemplo, foi construído em torno de 250 bilhões de páginas.

Esses 250 bilhões de páginas ou mais do que isso, porque diariamente são bilhões e bilhões de páginas que estão sendo incorporadas, de uma maneira grosseira, elas foram capturadas em um terço por organizações como governos, organizações não governamentais, empresas, entidades de uma forma geral - seguramente todo o conteúdo digital já produzido pelo Congresso brasileiro já foi ingerido por esses modelos -; um terço vem basicamente da produção acadêmica das universidades, enfim, da ciência; e um terço do jornalismo profissional, dos veículos de comunicação.

Só que ninguém sabe, no momento em que a gente faz uma busca, quais foram os conteúdos que foram capturados, usados e aí, diga-se com toda a transparência, sem nenhuma autorização. Esses conteúdos foram capturados, já estão ingeridos, já foram ingeridos pelos modelos de inteligência artificial sem autorização e sem qualquer pagamento, e isso vale também para os produtores de arte. O Dall-E não foi por acaso que faz desenhos e ilustrações; ele capturou milhões de ilustrações e desenhos e quadros e pinturas, enfim, e, a partir disso, também estabelece novos, supostamente novos quadros.

Então, nós não sabemos disso, e eu acho que isso é uma... E propositalmente não sabemos, porque já começam a acontecer nos Estados Unidos, na Europa, alguns processos de litigância, de cobrança de direitos autorais sobre o uso não autorizado de conteúdos jornalísticos, por exemplo, ou conteúdos artísticos.

Então, essa é uma coisa sobre a qual os governos e os Congressos, os Parlamentos têm que se debruçar: é sobre transparência, sobre direitos autorais, propriedade intelectual, porque senão nós vamos para um apocalipse informativo e para um apocalipse de direitos de autor.

Por exemplo, tem um *site*, uma organização chamada NewsGuard, que identificou que até junho - seguramente de junho para cá já tem muitos - já havia 150 supostos veículos de comunicação novos, pseudoveículos de comunicação, que são uma pessoa na sua garagem que captura conteúdos de dezenas de outros veículos, reembalha esses conteúdos, regurgita esses conteúdos segmentados ou não segmentados, como se fosse algo novo, sem nenhuma interferência humana, e captura apenas os cliques e, portanto, os anúncios, a receita correspondente a esses cliques, nesses *sites*, reescritos, reapresentados com conteúdos de terceiros.

Então, o desafio é gigantesco; assim, não existe uma solução única para todo esse emaranhado que nós estamos vendo de propriedade intelectual, de direitos, de respeito às leis estaduais, leis nacionais, mas uma coisa está clara: é preciso haver, sim, uma regulação - e não é uma questão de liberdade de expressão -, mas uma regulação no sentido de se estabelecer um mínimo de responsabilidade.

E, por fim, é a questão da desinformação. Os modelos de inteligência artificial não têm compromisso com o acerto, e nem tem ali também... Quando eles erram, pedem desculpa e refazem a informação. E outro: de quem é a responsabilidade legal por uma desinformação? De quem é a responsabilidade legal quando alguém repassa uma informação difamatória, por exemplo, extraída de um modelo de desenvolvimento de inteligência artificial? É do desenvolvedor? É de quem repassou? Quem repassou tem obrigação de saber que aquela informação é supostamente difamatória, ou que ela é caluniosa, ou que é injuriosa?

Então, esse desafio de quem tem a responsabilidade legal na origem seguramente está criando um novo dilema para todos os Parlamentos - e aqui também deixo essa contribuição.

Eu acredito, sim, que a ANJ endossa várias iniciativas nesse campo, como a Content Authenticity Initiative - que é a CAI -, que faz uma tentativa de estabelecer marcas d'água, identificações, talvez se tenha que usar *blockchain*... Aí eu deixo, obviamente, mais para os técnicos como é que a gente pode caracterizar que aquela informação tem certificado de origem.

Nós conseguimos fazer isso nos alimentos. Os alimentos têm um conteúdo X, os alimentos artificiais que estão naquela embalagem são identificados hoje na legislação brasileira, e talvez seja o caso também de nós procurarmos deixar claro para os usuários, para os consumidores de informação que aquilo tem um certificado de origem e quais são os ingredientes que compõem aquele certificado de origem.

Boa parte das redações está reescrevendo seus códigos de ética, mas isso é só o jornalismo profissional, que está - digamos assim - um pouco mais avançado nessa discussão. Paralelamente a isso, está surgindo uma série de iniciativas, como já mencionei, que não têm absolutamente nenhum compromisso com a informação profissional, com a ética, e isso, sim, é o que ameaça a estabilidade emocional e até geopolítica, num novo abalo, num novo *layer*, numa nova camada de desinformação a que nós começamos a assistir mundo afora.

Então, é muito louvável toda a preocupação, e contem com o nosso apoio para que a gente possa estabelecer uma regulação que seja uma regulação sensata, moderna e que, em nenhum momento, ameace o desenvolvimento responsável de uma tecnologia que tem que ser construída para o bem da humanidade, com o mínimo de efeitos negativos.

O Senador mencionou a energia nuclear, eu acho que é um excelente exemplo...

*(Soa a campainha.)*

**O SR. MARCELO ANTÔNIO RECH** *(Por videoconferência.)* - ... uma excelente analogia, e é isso que nós perseguimos também.

Então, bom dia a todos.

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado, Dr. Marcelo Rech, Presidente da Associação Nacional de Jornais. Parabéns pela apresentação. Eu fiz aqui umas anotações importantes, que eu vou levar para o Relator também.

Eu gostaria, neste momento, de passar a palavra para a Sra. Renata Mielli, Coordenadora do Comitê Gestor da Internet, por dez minutos.

Por favor, a senhora está remota, então, por favor, controle o tempo por aí.

Obrigado.

**A SRA. RENATA MIELLI** *(Para expor. Por videoconferência.)* - Obrigada.

Bom dia a todos e todas. Vocês me desculpem a voz meio estranha, mas eu estou com covid. Então, seguimos aqui.

Queria agradecer muito o convite, a iniciativa da Comissão especial, o convite promovido pelo Senador Eduardo Gomes, pelo Senador Astronauta Marcos Pontes, para a gente fazer este debate tão importante. Acho que os que me antecederam... Também é uma honra, um prazer estar aqui com tantos colegas com os quais eu venho trabalhando ao longo de muitos anos neste debate em torno do direito à comunicação, da liberdade de expressão, de uma internet que seja utilizada para o desenvolvimento da humanidade, para a solução de problemas. Acho que a gente está aqui diante de uma dimensão bastante desafiadora para discutir os impactos da inteligência artificial, no caso aqui, particularmente, nos processos eleitorais e no jornalismo.

Eu queria iniciar um pouco a partir de um certo óbvio, dizendo o que é uma inteligência artificial, que, na verdade, é um sistema automatizado, desenvolvido a partir de instruções e regras para executar uma tarefa específica. Nós ainda não temos a tal da inteligência artificial geral. Nós temos hoje sistemas que são desenvolvidos para executar tarefas específicas; e, para isso, esses modelos usam um grande volume de dados, cujo propósito é extrair padrões a partir de regras e instruções bastante definidas no momento do seu desenvolvimento. Portanto, nós estamos falando de sistemas estatísticos preditivos de alto poder de processamento de dados para definir padrões.

Quando a gente fala em definição de padrões, a gente pode pensar em muitas aplicações positivas. Extrair padrões climáticos para tentar enfrentar, por exemplo, ou evitar grandes catástrofes, a partir de padrões que você consegue identificar quando vai ter um evento climático extremo, é algo muito importante e uma aplicação muito nobre, eu diria, da inteligência artificial. Quando você consegue determinar eventos climáticos extremos e você pode, a partir disso, evitar a perda de um contingente imenso de produtos agrícolas, de alimentos, também você está usando a inteligência artificial para uma proteção econômica, social, alimentar da sociedade. Mas, quando você usa esse tipo de sistema para trabalhar com padrões comportamentais, nós podemos estar ingressando numa área bastante complexa e que pode trazer muitos danos para a sociedade como a gente conhece hoje.

Eu queria fazer um parêntese dizendo, eu sempre gosto de dizer isso, eu sou uma pessoa que tem dito isso muito, porque eu acho que tem a ver com o que nós estamos falando quando falamos de comportamento, de democracia e de jornalismo. Nós temos utilizado palavras extraídas de vocabulário usado para descrever reações humanas para nomear, para descrever esses sistemas tecnológicos maquínicos. A gente fala de inteligência, a gente fala de sistemas neurais, a gente fala de aprendizagem, a gente fala de decisões. Na verdade, essas palavras são, no mínimo, conceitualmente imprecisas e, portanto, bastante inadequadas para a gente lidar com essa tecnologia, principalmente no debate público, porque essas palavras criam uma mistificação em torno do que realmente essas ferramentas computacionais tratam e levam o usuário ao engano, que pode trazer muitos danos a partir dos seus usos, porque as pessoas que estão interagindo com essas tecnologias, a partir dessas denominações, acabam considerando que estão interagindo com uma tecnologia neutra, inteligente portanto, e a capacidade de questionamento dos usuários diante dos resultados que essas tecnologias nos oferecem é mínima. Então isso importa muito quando a gente está discutindo que tipo não só de regulação, mas qual o tratamento que a sociedade vai dar para empoderar o cidadão diante dessa tecnologia.

Como disse a Patricia, nós estamos diante de um sistema em que, primeiro, para poder fazer todo esse cruzamento a partir de milhões de dados, esses dados, e não só os dados, mas desde o desenvolvimento desse sistema, passando pela programação e o treinamento com base nesses dados volumosos, carregam vieses de origem, e esses vieses são os mais variados possíveis. E não estamos falando apenas de vieses, estamos falando também do propósito para o qual aquela determinada solução ou tecnologia de inteligência artificial está sendo usada. Então, quando a gente discute isto, sobre os impactos eleitorais e democráticos, e mesmo no jornalismo, a gente se coloca diante de desafios e preocupações que especialistas têm apontado no geral. E muitas das preocupações já vieram, já compareceram aqui na nossa audiência, não é? E, no desenvolvimento da IA - e se a gente for pensar na IA como uma ferramenta tecnológica -, digamos que os estágios anteriores dela foram aqueles estágios utilizados, por exemplo, para direcionar conteúdo individualizado no escândalo do Cambridge Analytica. Aquilo era um algoritmo utilizado para distribuir conteúdo, mensagem eleitoral, com o propósito de persuadir o eleitor a partir de um perfilamento emocional, comportamental daquele usuário. Nós evoluímos nesse modelo de linguagem, nesse modelo de inteligência artificial, cuja capacidade de extrair, com ainda mais precisão, o padrão eleitoral e político é imensa, não é? E, por isso, a gente precisa procurar refletir sobre esses impactos.

A primeira questão já compareceu aqui. Nós estamos falando de um sistema que tem completa ausência de transparência. Nós não sabemos quais os mecanismos utilizados pelo modelo de processamento de linguagem para chegar a determinado resultado, nem sabemos qual é o conjunto de dados que foram utilizados para oferecer aquele resultado. E aquele resultado é um resultado personalizado. Muito provavelmente, se eu fizer uma pergunta para o ChatGPT agora, e qualquer outra

pessoa desta sala fizer a mesma pergunta, é possível que o resultado dessa resposta seja bastante diferente, de acordo com o perfil que é traçado para mim.

Então essa situação nos leva a uma ruptura no sistema de confiança entre sociedade e informação, porque, se nós já estamos passando por um processo de deterioração da esfera pública e de perda de referenciais com credibilidade, esse tipo de avanço da tecnologia de inteligência artificial e de modelos de linguagem para oferecer resultados coloca ainda mais um alerta sobre essa questão de confiabilidade da informação. Isso, porque esses sistemas conseguem gerar mensagens diferentes sobre o mesmo assunto para milhões de pessoas, não é? E ela vai, no caso, ao ser aplicada em processos eleitorais, ter uma tarefa, e essa tarefa é buscar dissuadir e influenciar o voto daquela determinada pessoa.

Então não importa o conteúdo, importa o resultado, porque essas ferramentas não têm uma ética, não é? Por mais que a gente pode pensar numa ética *by design*, mas elas não têm moral. Elas não têm compromisso com a verdade necessariamente. Então isso é um desafio imenso.

E, como disse o Prof. Diogo Rais aqui, elas enviam as mensagens que nós gostaríamos de receber sobre candidatos que nós gostaríamos de conhecer, porque temos afinidades com as suas ideias e, portanto, temos uma predisposição para votar neles, se a gente for pensar num processo eleitoral. Isso, em parte, é positivo.

De outra parte, é bastante negativo, porque, se a gente pensa no ideal de um sistema democrático, uma democracia de alta intensidade precisa contar com informações com as quais concordo, mas eu preciso ser exposta às informações que a confrontam, porque eu preciso estar diante da opinião do outro para chegar a uma compreensão e a uma visão mais crítica da sociedade. Então, isso é um outro desafio.

Meu tempo está acabando.

Eu fiz uma experiência: eu peguei o nome e o tema da nossa audiência, coloquei no ChatGPT e pedi para ele me devolver algumas ideias sobre esse assunto. Entre as ideias que ele me ofereceu sobre o assunto, é exatamente a de como a inteligência artificial, no estágio em que nós estamos hoje, pode aprofundar a polarização em função do direcionamento de mensagens personalizadas que colocam para o usuário apenas aquilo que ele gostaria de ouvir. E, para mudar a opinião das pessoas, uma das principais estratégias de mobilização no *marketing* político, aplicado a partir da inteligência artificial, é a aplicação das emoções, porque são as emoções as estratégias mais rápidas e eficientes para alterar e reforçar crenças a partir do medo, da raiva.

Portanto, eu acho que nós temos um grande desafio.

Meu tempo acabou.

Eu queria acrescentar uma última questão bem rapidamente. Eu acho que, além de rotular os conteúdos de IA, nós precisamos pensar a partir de um outro paradigma. Não basta o usuário saber que ele está interagindo com um conteúdo ou com uma mensagem que foi produzida por IA, nós precisamos pensar diferente: qual o objetivo daquele conteúdo, se ele foi um conteúdo produzido por uma propaganda eleitoral, quem produziu aquele conteúdo e quem distribuiu aquele conteúdo. Eu acho que esse é um conjunto de informações essenciais para que a gente possa fornecer mais empoderamento. Não basta dizer: conteúdo produzido por IA, mas dizer: foi produzido para quê, com qual propósito, por quem e por quem está sendo distribuído.

Então, essas seriam algumas das minhas contribuições.

Agradeço aí.

Desculpe-me por ter passado um pouquinho do tempo.

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado, Sra. Renata Mielli, Coordenadora do Comitê Gestor da Internet. Parabéns pela apresentação e votos de melhoras aí também.

Mas, realmente, pontos muito importantes, e aqui eu tomei nota de vários deles.

Sem dúvida nenhuma, a gente tem que pensar, e o que sempre vem à mente é o desenvolvimento do ser humano em paralelo a isso. Então, eu estava citando aqui... Eu faço parte da Comissão de Educação e falo muito de educação, mas a educação - parece clichê - está na base de tudo isso. A gente precisa realmente ter essa educação, não só a educação formal, dos conteúdos formais, mas também a educação no sentido intrapessoal, do indivíduo, de conhecer.

Então, parabéns pela apresentação!

Neste momento, então, eu passo a palavra para a Sra. Tainá Aguiar Junquilho, Professora de Direito, Inovações e Tecnologia do Instituto Brasileiro de Ensino, Desenvolvimento e Pesquisa - está aí o ensino.

V. Sa. tem dez minutos.

Obrigado.

**A SRA. TAINÁ AGUIAR JUNQUILHO** (Para expor.) - Bom dia a todos e a todas.

Eu queria agradecer ao Senador Eduardo Gomes pelo convite e ao Senador Astronauta Marcos Pontes.

Parabenizo esta Casa também pela iniciativa, tanto da criação desta Comissão quanto das audiências realizadas aqui, para a gente lapidar o texto que foi proposto pela Comissão de Juristas, no ano passado.

Eu confesso que é bonito de ver os senhores anotando e interagindo com os interessados no tema também pelo e-Cidadania.

Eu faço parte da Coalizão Direitos na Rede, que emitiu uma nota sobre o PL que eu vou entregar ao final aqui da minha apresentação. Sou advogada e faço parte também da Comissão Nacional da OAB sobre direito digital, e pesquisadora do IDP, exatamente em regulação de inteligência artificial. Então, eu me sinto muito honrada de estar aqui na manhã de hoje colaborando com essas discussões.

E a minha fala aqui hoje está dividida em três pontos. Na primeira parte, coube-me aqui discutir um tema que eu acredito ser muito sensível para esta Casa, que é inteligência artificial e democracia. E aí, em relação aos pilares democráticos, são várias as afetações pelos sistemas de inteligência artificial. Então, em primeiro lugar, a inteligência artificial tem tido muitos usos no setor público, por exemplo, para ingresso de pessoas na imigração de países, nas alfândegas. O próprio INSS, aqui no Brasil, para concessão ou não de benefícios. O Poder Judiciário brasileiro já conta com 44 tribunais dos 92 do país utilizando e desenvolvendo algum tipo de projeto de inteligência artificial. Eu vi uma notícia recente também de que escolas no Paraná estão utilizando inteligência artificial para análise de sentimentos e emoções. E essas aplicações precisam de transparência para o regime democrático. Isto é, o cidadão precisa compreender por que e como essas políticas públicas estão sendo ou não oferecidas, denegadas, concedidas, e no que elas o afetam efetivamente.

O sistema democrático também tem recebido diversos avisos sobre o poder algorítmico em eleições. Em 2016, a gente teve o escândalo da Cambridge Analytica, que já foi citado aqui na manhã de hoje. Um estudo do MIT, inclusive, chamado "Quantificando Potenciais Retornos Persuasivos da Microsegmentação Política", descobriu que os algoritmos da Cambridge Analytica foram 70% eficazes para influenciar comportamentos e votos e produzir conteúdos personalizados.

Eu não sou tão antiga assim, mas sou da época em que um político, para se eleger, contratava um marqueteiro para criar um bom *jingle* - eu me lembro de alguns ainda. E hoje não é assim mais, com a inteligência artificial. A gente tem também os efeitos filtro-bolha e câmaras de eco, descritos por Eli Pariser, que são uma filtragem algorítmica que produz impulsionamento e nos coloca em verdadeiras bolhas de preferência, aumentando também a polarização pelo mundo, e que acabaram influenciando em eventos de ataques importantes à democracia como o 6 de janeiro nos Estados Unidos, com a invasão do Capitólio, e o 8 de janeiro aqui no Brasil.

Em 2022, também, lançamentos da inteligência artificial generativa, que é essa que cria áudios, textos, imagens e conteúdos que podem ser desinformativos e enganar os eleitores, e aí você tem inteligência artificial já criando imagens de Trump sendo preso; em eleições na Coreia do Sul, os candidatos utilizaram avatares - o Diogo já comentou isso - na campanha; na eleição nacional da Eslováquia; o Reino Unido também já utilizou para fazer uma *deepfake* contra o candidato a Primeiro-Ministro.

E, no Brasil, dados do TSE nos dão conta de que o comparecimento médio de jovens de 16 a 17 anos aumentou 52,3% entre 2018 e 2022. A tendência é que isso cresça cada vez mais. Os jovens estão na internet, eles estão em aplicativos que vão coordenar a nossa atenção. Então, no ano que vem, principalmente em 2024, nas próximas eleições, a gente vai ter desafios gigantescos. São 5.070 municípios para os quais o Estado vai precisar estar preparado para responder a esses desafios da inteligência artificial. "Ah, mas nós não vamos estar preparados para regular". Olha, nós temos que, enquanto Estado, dar uma resposta. A integridade da democracia está sob ameaça, e o relógio está contando para nós tomarmos alguma atitude. Não podemos deixar para ver no que vai dar.

Na segunda parte, quanto aos debates até aqui realizados em audiência, eu queria fazer dois destaques. Foi criticada a existência da agência reguladora porque inteligência artificial é uma tecnologia de propósito geral, e isso exigiria conhecimentos setoriais. Ora, gente, se a gente cria uma regulação sem estabelecer uma agência fiscalizadora, essa regulação vai estar fadada ao fracasso, porque o controle disso vai ficar granularizado. Quem vai fiscalizar? Quem vai auditar?

Quando surgiu a internet, naquela época, surgiu também uma forma de governança muito interessante, que o Comitê Gestor da Internet faz, que é a governança com a multissetorialidade, aquela que envolve vários setores. Agora, com a tecnologia da inteligência artificial, que é tão disruptiva, a gente também pode propor um órgão fiscalizador que tenha não só multissetorialidade - vários setores no Brasil representados -, mas também multidisciplinaridade, um órgão em que

cada agência tenha o seu assento, ou, então, órgãos reguladores setoriais que atuem de forma coordenada com esse órgão regulador central, como é na União Europeia e agora também nos Estados Unidos.

Esse órgão central é fundamental para identificar riscos, para dar segurança jurídica para o desenvolvimento da inteligência artificial, para promover políticas públicas de educação, *hackathons* de inteligência artificial para o Brasil, promover, como já define o art. 32 do PL 2.338, a implementação da estratégia brasileira de inteligência artificial

Em segundo lugar, outro ponto das audiências que eu queria destacar é abordar a falácia que foi aqui levantada, em algumas falas, de que a regulação impede a inovação. A falácia, gente, é a construção de uma ideia que aparenta uma lógica de verdade, mas possui um ponto que a torna falsa. E aqui esse dilema falso, inovação *versus* regulação, é a falácia do falso dilema, que vai apresentar ali duas alternativas só, quando, na verdade, a complexidade da realidade é muito mais ampla. A gente precisa, sim, regular e ir além de uma principiologia que já existe, colocada por órgãos internacionais, como, por exemplo, a OCDE, a Unesco, e existem várias possibilidades além disto de você incluir medidas que apoiem o desenvolvimento de uma IA responsável numa regulação.

De mais a mais, a inovação exige uma segurança, e eu não estou sozinha nisso, não é? Luciano Floridi, no último livro dele, *The Green and the Blue*, fala dessa conciliação entre o verde da sustentabilidade e o azul da tecnologia. A Virginia Dignum também fala da responsabilidade, e, no meu livro também, sobre ética de inteligência artificial, eu proponho - aqui fica a dica para os Senadores - uma PEC para incluir, no rol dos direitos e garantias fundamentais, a produção de uma inteligência artificial responsável como direito fundamental.

A terceira parte, para finalizar, seria o que a gente pode ajustar no PL 2.338. O primeiro ponto que eu queria deixar aqui assentado é consenso. A gente precisa ir além da principiologia que já existe, e isso não sou só eu que estou falando; são os pesquisadores, são os próprios CEOs das empresas que estão pedindo regulação. A União Europeia já tem, já está em vias de aprovação do AI Act, e ontem, como foi aqui falado também na manhã de hoje, os Estados Unidos soltaram uma ordem executiva requerendo das agências setoriais que produzam relatórios e constituam os chamados *red teams*, ou times vermelhos, que vão adotar e entender como está sendo a adoção de uma série de medidas para controlar os riscos da inteligência artificial dentro de cada setor.

O Brasil como fica? Ele tem um potencial para produzir inteligência artificial para diagnósticos médicos, para vacina, na agricultura... Nós somos um país gigantesco de língua lusófona portuguesa, então que pode despontar inteligência artificial generativa na língua portuguesa.

(*Soa a campanha.*)

**A SRA. TAINÁ AGUIAR JUNQUILHO** - Nós somos um povo diverso em imagens também. Então, isso ajuda na produção de inteligência artificial para imagens.

E o PL, então, pode alterar as definições de ciclo de vida da IA para incluir conceitos de desenvolvedor, fornecedor e aplicador. Além das *sandboxes*, inserir incentivos também para estimular as *startups* brasileiras, como fez a União Europeia no desenvolvimento de inteligência artificial responsável, prever investimentos em educação e literacia digital, além de políticas públicas de apoio a incentivos a quem promover a inteligência artificial responsável - isso poderia caber a um órgão central -, previsão dessa autoridade, como fizeram os Estados Unidos ontem, inclusive, demandando de cada agência que faça reiterados relatórios de riscos, para subsidiar a atuação dessa autoridade.

Em relação à inteligência artificial generativa, que foi lançada uma semana antes da entrega do relatório pela Comissão...

(*Soa a campanha.*)

**A SRA. TAINÁ AGUIAR JUNQUILHO** - ... e que, portanto, não tem nada sobre ela ainda no PL 2.338, recomendamos aqui o selo de identificação do conteúdo, o estímulo de empresas de IA generativa para a língua portuguesa no Brasil e também a pesquisas.

Bom, para finalizar então, que o meu tempo já passou - a vizinha dos 15 segundos já soou -, esta Casa tem diante de si, Senador, uma decisão histórica. E qual mensagem os senhores querem passar?

Eu tenho certeza de que os legisladores desta Casa querem propor uma regulamentação que preserve a integridade da nossa democracia e que, portanto, estimule e assegure a soberania do Brasil e o desenvolvimento dessa tecnologia de forma responsável no país.

O Brasil conta com os senhores, e eu também me coloco à disposição para ajudá-los no que for necessário nesse tema que é tão difícil, mas importante, para o nosso país.

Obrigada.

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado, obrigado à Sra. Tainá Aguiar Junquillo, que é Professora de Direito, Inovação e Tecnologia do Instituto Brasileiro de Ensino, Desenvolvimento e Pesquisa.

Tomei várias notas aqui a respeito do assunto que foi tratado e como foi tratado. Inclusive, algumas vezes bastante interessante, com um sentido à parte internacional.

Eu era o representante do país no Gpai, no Global Partnership for Artificial Intelligence, e analisamos lá várias das legislações que estão sendo aplicadas ou desenvolvidas ainda em vários países - e algumas já com modificações -, e, sem dúvida, a gente precisa garantir que a nossa aqui utilize esse conhecimento já desenvolvido em outros países e também um ponto importantíssimo, que nos garanta essa soberania, porque esse é um ponto que não tem sido trazido muito aqui, mas, sem dúvida nenhuma, perder mão nesse sentido afeta diretamente a soberania.

Parabéns, Tainá. Obrigado.

Gostaria de então passar a palavra, neste momento, para a Profa. Dora Kaufman, Professora da Pontifícia Universidade Católica de São Paulo (PUC-SP).

Professora, a senhora tem os dez minutos. Peço para controlar pelo lado de lá.

Obrigado.

**A SRA. DORA KAUFMAN** (Para expor. *Por videoconferência.*) - Bom dia.

Cumprimento a todos os Senadores e todos os meus colegas.

Agradeço o convite da Comissão do Senado e, especialmente, a gentileza de lidar com a minha agenda, o que possibilitou garantir a minha oportunidade de contribuir com os Senadores e com o país. Muito obrigada por isso.

O tempo é curto, e o tema é complexo. A vantagem de falar entre os últimos é que os colegas que me antecederam já abordaram diretamente alguns dos temas da audiência de hoje, o que me permite fazer contribuições mais gerais sobre o processo.

A IA é estratégica para o desenvolvimento do Brasil e das organizações públicas e privadas, eu acho que isso é um consenso entre nós todos. A IA é crítica, inclusive, para a gestão pública. Eu acho que grande parte da eficiência da gestão pública poderia ser alcançada com o uso da inteligência artificial. Nas várias áreas do poder público, eu tenho a impressão de que hoje, proporcionalmente, mais áreas do poder público estão usando a inteligência artificial no Brasil do que mesmo o setor privado. A competitividade nacional do país e das empresas nacionais, as nossas empresas, depende também de criar condições favoráveis ao desenvolvimento e uso da IA. Junto com qualquer regulação, que eu sou a favor, é preciso criar infraestrutura adequada, recursos para pesquisa e normas, senão a gente não avança.

Dada a imensa relevância da IA, eu acho que não podemos nos precipitar. Vale observar que não existe um marco regulatório da IA em nenhuma região do Ocidente neste momento. E esse fato não é por acaso, reflete a dificuldade do tema. O processo europeu, que é o mais avançado e o mais robusto, teve início em 2018. Em 21 de abril de 2021, foi lançada para a consulta pública a primeira versão, que, em novembro do ano passado, já tinha 3 mil emendas. Em 14 de junho deste ano, teve uma primeira votação no Parlamento europeu, que fez uma série de modificações, que, no meu ponto de vista, tornou a proposta regulatória europeia muito mais complexa - para mim, eu tenho várias dúvidas sobre a capacidade de viabilização dessa versão que está hoje, a última versão. O marco relatório da internet, também para a gente ter uma referência nacional, demorou cinco anos em debate, e esse marco é uma referência no mundo todo.

O PL 2.338 data de 3 de maio. Então, nós estamos falando de somente seis meses em ele se transformou num PL. É muito pouco tempo para qualquer decisão definitiva que vai ter um impacto tão grande no futuro do país.

Semanalmente, também a gente está acompanhando, surgem inúmeras contribuições que também podem nos servir de referência. Por exemplo, ontem, o que já foi citado, o Governo americano, o Governo do Presidente Biden, divulgou um documento, uma ordem executiva com dois focos principais: segurança e proteção, infraestrutura de pesquisa e implementação financiada pelos recursos públicos. Esse documento não tem poder de lei, apenas são diretrizes, coerente com o posicionamento da Casa Branca, que evolui baseado em legislação e estruturas existentes, o que é distinto da Europa, mas são iniciativas que podem nos dar referências, são iniciativas nos Estados Unidos e na Europa simultaneamente que podem nos dar referências.

Outra iniciativa que eu gostaria de comentar é que, semana passada, a ONU constituiu um órgão consultivo de alto nível sobre inteligência artificial. Esse órgão é composto de 32 especialistas - não necessariamente especialistas em inteligência oficial, alguns nem são -, e o propósito é promover uma abordagem globalmente inclusiva. Pelo Brasil, nós temos a Assessora Especial do Ministro da Justiça e Segurança Pública do Governo Federal Estela Aranha. Então, são duas -

uma aconteceu ontem, e outra aconteceu semana passada - iniciativas importantes que agregam nesse debate sobre a inteligência artificial no âmbito global.

Voltando ao PL 2.338, eu já falei isso em vários fóruns, eu o considero um ponto de partida, sem dúvida é o melhor projeto de lei em tramitação no Congresso sobre inteligência artificial; os demais me parecem inócuos, tanto no sentido de proteger a sociedade quanto de estimular a inovação. Mas também, na minha opinião, o PL 2.338 não está pronto para se tornar um marco regulatório da IA no Brasil.

No dia 11 de agosto eu publiquei, no jornal *Valor Econômico*, um artigo analisando detalhadamente o PL 2.338. Aqui, dado o curto tempo, eu vou só fazer essa menção de que o artigo está lá com as minhas opiniões.

Os modelos de IA generativa, que todos nós sabemos que é o fenômeno deste ano, não estão contemplados, por exemplo, no 2.338. E esse modelo, essas soluções de IA generativas acrescentam um novo grau de complexidade que precisa ser investigado. Existem inúmeras questões não contempladas ou não adequadamente contempladas pelo PL 2.338. Algumas já foram citadas pelos colegas e muitas outras não foram ainda citadas.

Além disso, a eficiência da regulamentação depende de se definirem padrões, normas, procedimentos de arbitragem, acordos internacionais, indicadores que serão investigados e avaliados, instrumentos de fiscalização. São um enorme desafio os instrumentos de fiscalização, como fiscalizar, porque não adianta também ter um marco regulatório se não tem instrumento de fiscalização. Então como fiscalizar uma tecnologia pouco ou nada transparente é uma complexidade à parte. Quem vai e como será feita essa fiscalização? Uma das motivações, inclusive, para minha defesa de o protagonismo ser das agências setoriais é exatamente isso. O que eu acho é que a pertinência da regulação específica e da fiscalização deve caber a essas agências setoriais. Estou defendendo isso desde o início desse processo de discussão aqui no Brasil, sem descartar uma espécie de agência central a ser definida como será a composição e atribuições, uma espécie de coordenação que eu ainda não tenho claro. Mas me parece ser o único caminho de efetivamente termos uma regulamentação sendo implementada e fiscalizada é o protagonismo ser das agências reguladoras setoriais.

Minha preocupação no momento é aprovarmos algo tão complexo que tenha pouca viabilidade, favorecendo a concentração de mercado, porque sempre as grandes empresas têm mecanismos de responder a qualquer que seja a regulamentação vigente, ou uma regulamentação tão generalista que seja inócua na missão de proteger a sociedade. Em suma, o que eu queria transmitir é que eu não creio que o tema da regulamentação já esteja suficientemente maduro no Brasil para se estabelecer neste momento o marco regulatório para o desenvolvimento e o uso de uma tecnologia estratégia para o nosso país. Apesar de eu ser absolutamente a favor de existir um marco regulatório, acho só que não estamos preparados para isso.

Essa constatação deriva de inúmeros fóruns de que tenho participado no Brasil e no exterior.

No momento, eu preciso desligar aqui - só um minuto, por favor - que está tocando o telefone. *(Pausa.)*

Desculpe-me. Deve ter uns quatro anos que não toca o telefone fixo na minha casa; tocou justamente neste momento, hoje.

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Existe ainda?

**A SRA. DORA KAUFMAN** *(Por videoconferência.)* - Existe, mas eu nem uso nem toca, mas hoje tocou.

Então, essa constatação de que o PL 2.338 não está pronto para se tornar o marco regulatório da inteligência artificial no Brasil, não está maduro e a sociedade brasileira não está madura ainda para definir o que seria o teor desse marco regulatório deriva de vários fóruns. Eu tenho participado intensamente dos fóruns mais distintos no Brasil e no exterior. De fato, hoje, também o depoimento dos colegas - com todo o respeito e carinho - mostra quanto a gente ainda não está preparado, enquanto sociedade, para estabelecer esse marco regulatório, que vai ser fundamental para o futuro do país.

Então, eu tenho uma sugestão objetiva.

Eu tenho, como sugestão para a Comissão do Senado, que seja formado um conselho composto de representantes das agências reguladoras setoriais, de representantes de institutos de pesquisa das universidades, constituindo uma secretaria-executiva rotativa, porque é sempre importante ter uma secretaria-executiva, com a missão inicial de propor um arcabouço regulatório, contemplando os desafios de cada domínio.

Então, claro, que o início do debate desse conselho, desse fórum que eu estou sugerindo seja o PL 2.338, com o olhar de quais são as implicações e os desafios em cada um dos domínios de aplicação.

Posteriormente, esse conselho poderia ser responsável por coordenar as ações regulatórias fiscalizadoras, ou seja, esse conselho poderia se transformar na agência coordenadora, agência central, mas isso é uma discussão depois.

Então, esse conselho eu vejo como um órgão fundamental, como uma etapa fundamental.

Há uma diferença. Eu aprecio e reconheço todo o esforço dos Senadores. Sinceramente, acho mesmo um grande esforço fazer esse processo, organizar essas consultas públicas, mas todos nós temos consciência de que o tempo é muito curto. Dez minutos para abordar um assunto tão complexo restringe a nossa contribuição. E mesmo enviar contribuições por escrito, todas essas iniciativas que são válidas, tudo isso é limitado. É muito distinto de formar...

Eu lembro que o processo de regulamentação da Europa, por exemplo, que tem vários problemas que não estão equacionados, na minha opinião, começou com... Em 2018, começou com a formação de um observatório que levantou como estava sendo usada a inteligência artificial na Europa, o que, pelo que eu saiba, a gente não fez sistematicamente aqui no Brasil, e criou-se uma comissão de - se não me engano, eram 56 especialistas, de várias áreas, que produziam os primeiros documentos...

Então, eu acho que a gente não tem tempo para voltar atrás, mas a minha sugestão é esta: que a Comissão do Senado - o Senado - dê urgência à discussão do marco regulatório, mas não dê urgência para aprovação de um arcabouço regulatório de inteligência artificial no Brasil. Que se forme esse conselho, chamem-se as agências reguladoras setoriais e os institutos de pesquisa das universidades, principalmente - e outras organizações que se julgarem necessárias, importantes, estratégicas -, para discutir, para formar um grupo de trabalho, coordenado por uma secretaria...

Então, é isso, são essas as minhas contribuições.

Eu acho que esses órgãos reguladores setoriais serão capazes, inclusive, de estabelecer fóruns com o mercado, com as empresas dos seus domínios, trazendo as contribuições desse segmento, quer dizer, o segmento do setor privado, para o conselho.

Então, em suma, a minha contribuição é mais no sentido de uma sugestão de processo do que do conteúdo em si da 2.338, que, como eu falei, publiquei já no *Valor Econômico* e não tenho nada a acrescentar desde lá.

Muito obrigada.

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado. Obrigado, Profa. Dora Kaufman, Professora da Pontifícia Universidade Católica de São Paulo (PUC-SP).

Eu anotei aqui diversos pontos importantes também e estava até comentando, aqui de lado, sobre essa necessidade da participação das agências reguladoras. Nós estivemos aqui já, nessas audiências, com a saúde; nós tivemos audiência com o setor de proteção de dados; nós tivemos audiência trazendo também a Anatel e outros. Então, é importante realmente ter esse ponto de vista que vai por várias perspectivas.

Finalmente, eu gostaria de passar a palavra ao nosso último apresentador, que é o Sr. João Caldeira Brant Monteiro de Castro, Secretário de Políticas Digitais da Secretaria de Comunicação Social da Presidência da República.

Sr. João Brant, o senhor tem dez minutos.

Por favor, peço para controlar o tempo por aí, porque a gente não tem como mostrar aqui.

Obrigado.

**O SR. JOÃO CALDEIRA BRANT MONTEIRO DE CASTRO** (Para expor. *Por videoconferência.*) - Está bem.

Muito obrigado, Senador Astronauta Marcos Pontes. Agradeço também ao Senador Eduardo Gomes pelo convite, pela possibilidade de trazer as visões aqui da Secretaria de Políticas Digitais.

A secretaria foi formada neste ano, justamente com a missão de olhar para o que a gente chama de interação entre o ambiente informacional e o ambiente digital. Portanto, é dedicada a discutir os efeitos da desinformação, o impacto do ambiente digital na sustentabilidade do jornalismo, a questão da educação midiática, uma série de temas que já foram trazidos aqui e que são objeto desta audiência de hoje.

Eu vou começar dando um pouco o lide da minha fala e acho que ela vai se justificar depois, mas eu acho que nós precisamos olhar para esse marco regulatório de inteligência artificial como um marco que...

Vou só esperar, porque é o lide da minha fala e eu queria poder dialogar diretamente consigo também, Senador.

Eu acho que o marco precisa dar conta de uma definição mais geral e de proteger, digamos assim, lacunas existentes nas políticas setoriais, mas ele não deve passar por cima de respostas setoriais aos temas e às implicações de inteligência artificial. Acho que isso é essencial hoje para a gente olhar pensando no tema de integridade da informação e no tema eleitoral.

Nós temos tratado a discussão da desinformação, do discurso de ódio e do enfrentamento a esses temas como resposta a uma busca pela integridade da informação. Esse é um tema que as Nações Unidas, o Pnud, a OCDE e outros atores têm adotado como uma agenda, vamos chamar assim, positiva do que está se tentando proteger. A ONU define isso a partir de uma definição de confiabilidade, precisão e consistência das informações. Portanto, isso está ligado à ideia de uma,

digamos, dimensão coletiva da liberdade de expressão, da ideia de que é preciso proteger o direito de a sociedade estar bem informada para sua tomada de decisões.

E me parece que a defesa da integridade da informação precisa ser um ponto central ao se olhar a discussão de inteligência artificial, porque ela compõe um ponto de partida, ela deveria estar sendo inclusive trazida como um dos princípios organizadores, porque ela compõe um ponto de partida chave aqui no que a gente está buscando proteger.

Eu acho que já foi extensamente discutido e percorrido pelos meus colegas sobre os efeitos positivos e negativos da inteligência artificial. É evidente que, quando a gente olha para um marco regulatório, nós estamos tentando preservar a metade cheia do copo e afetar ou proteger os cidadãos em relação à metade vazia, aos efeitos negativos, e é por isso que eles acabam sendo mais enfatizados aqui nesse espaço. Mas nós temos que ter a mesma preocupação de não esvaziar a metade cheia do copo, para que a gente não acabe, a partir de um marco sobre inteligência artificial, afetando os efeitos positivos que ela tem trazido.

Eu acho que já a Relatora para a Liberdade de Expressão e Opinião, da ONU, vem discutindo como é que determinados sistemas de informação, ao provocarem uma organização de algoritmos baseada na ideia de maximização de engajamento, desorganizam a esfera pública de debate e fazem com que você afete a liberdade de opinião das pessoas ao impor padrões de comportamento e preferências ocultas nessas interações nas redes sociais. Então, acho que nós não podemos perder de perspectiva que parte dos efeitos negativos que a gente está discutindo aqui são fruto justamente da modelagem de algoritmos baseada no objetivo de maximização do engajamento do usuário como marco principal de organização desse ambiente informacional. O que eu quero dizer com isso, trocando em miúdos, é: tudo que a gente fez e aprendeu no fortalecimento do jornalismo profissional entre o pós-guerra e 2009 e 2010, nós desorganizamos na organização do ambiente informacional nos últimos 10, 15 anos, e acho que isso não pode ser perdido de consideração.

Eu acho que o formato do PL 2.338 está adequado ao tamanho do desafio, a ideia de você lidar com transparência e devido processo, de você trazer a discussão de responsabilidade civil organizada condicionada ao tamanho do risco que ele carrega e de você trazer avaliação e atenuação de riscos sistêmicos. Inclusive, é uma abordagem muito parecida com o que o Governo tem defendido no PL 2.630, de origem do Senado, que está em debate na Câmara e que trata da regulação do ambiente informacional nas redes sociais. Parece-nos que essa é uma abordagem que precisa ser considerada quando a gente olha para o tema da inteligência artificial. Quando a gente olha a inteligência artificial afetando o ambiente informacional, nós estamos falando de temas muito próximos aos que estão em discussão lá no PL 2.630, e acho que, portanto, essas normativas precisam claramente conversar e, de certa maneira, se dizer o que se impõe ao quê. E aí eu repito o que eu falei no começo: acho que as normas específicas, como as do PL 2.630 no caso das redes sociais, se impõem sobre uma norma mais geral de abordagem da inteligência artificial, mas o ideal é que a gente tenha isso de maneira clara e conversada, para que esses normativos não se choquem. Acho que, no caso do 2.630, nós temos defendido também a ideia de um dever de cuidado das plataformas na proteção contra a disseminação de conteúdos ilegais. Acho que esse elemento pode inspirar também uma parte do debate que se faz aqui quanto ao PL sobre inteligência artificial.

O impacto eleitoral nos parece bastante relevante também de ser trazido para a primeira tona, mas, de novo: não sei se é o PL 2.338 que deve tentar resolver as discussões de impacto eleitoral. Ele precisa talvez endereçar essas questões para a Lei Geral de Eleições, para o Código Eleitoral, para as questões que estão sendo tratadas ali. Nós vemos a inteligência artificial se transformar numa ferramenta de manipulação cada vez mais comum, candidatos estão recorrendo à IA para criar chamadas, mensagens e imagens que por um lado podem facilitar um processo de propaganda eleitoral, mas por outro podem distorcer a realidade, e aí é crucial que esses atores assumam esse compromisso ético. Com IA, eles devem incluir obviamente rotulagem de qualquer comunicação gerada por ferramentas de IA, proibição de uso da IA para difamar concorrentes ou enganar eleitores, bem como evitar o uso da IA para confundir as pessoas sobre como votar.

Mas também não se limita a uma discussão de candidatos específicos, nós estamos falando da possibilidade de corroer a confiança dos cidadãos na integridade do processo democrático como um todo; e a facilidade com que a IA pode criar conteúdo falso, a gente sabe que é impressionante. Portanto, é essencial que os candidatos revelem quando estão usando IA para criar comunicações, mas também que toda a disseminação de informação de IA que possa causar um efeito enganador no eleitor, seja por edição de autoria ou por edição dos próprios elementos fáticos que ela tenta reproduzir, estejam claramente evitadas e cercadas pelo normativo eleitoral.

E é crucial também que, portanto, *deepfakes* perigosos, que mostrem candidatos fazendo ou dizendo coisas que nunca fizeram, não sejam permitidos. E aqui especialmente no período eleitoral, nós precisamos trabalhar esse elemento.

Acho que... Terminando dizendo o seguinte: acho que a garantia que a gente precisa dar... Bom, desculpa, tem um elemento a mais que eu queria observar, reforçando uma observação que o Marcelo Rech colocou, que é em relação à possibilidade de o desenho de direito autoral previsto na lei afetar a sustentabilidade do jornalismo no ambiente digital. Portanto, a gente

precisa olhar, rever e buscar proteger os direitos patrimoniais dos produtores de jornalismo do ambiente digital, para que a gente não tenha uma exploração indevida e na prática uma afetação negativa na sustentabilidade financeira.

Então, o que eu queria ressaltar aqui é que as normas que estão colocadas hoje e que estão sendo discutidas em relação a redes sociais, que estão colocadas nos regulamentos de serviços de radiodifusão, que estão colocadas no Código Eleitoral e que precisam ser discutidas aqui...

*(Soa a campanha.)*

**O SR. JOÃO CALDEIRA BRANT MONTEIRO DE CASTRO** *(Por videoconferência.)* - ...em relação ao efeito eleitoral, elas precisam ter força específica.

Portanto, nesse sentido, eu me alio à fala da Dora Kaufman, quando acha que boa parte do que a gente precisa fazer é fortalecer a regulação específica e setorial no tratamento desses temas.

Então, essa discussão sobre o equilíbrio entre que em medida a norma de inteligência artificial vai entrar cobrindo lacunas e estabelecendo proteções gerais? Em que medida ela precisa estar subjugada a regulamentos específicos? Parece-me um tema central sobre o qual o Governo se propõe a discutir, até porque isso impacta na discussão sobre autoridade, como já expressou a própria Dora aqui, e onde o Poder Executivo encontra...

*(Soa a campanha.)*

**O SR. JOÃO CALDEIRA BRANT MONTEIRO DE CASTRO** *(Por videoconferência.)* - ... prerrogativa de definir e de incidir.

Então, é isso. Eu agradeço. E agradeço a presença da Samara Castro, que trabalha comigo, é Diretora de Promoção da Liberdade de Expressão, que está presente aí no Plenário e também pode ajudar a destrinchar um pouco isso, nessa pós-audiência aqui.

Nós já estamos em contato também com o gabinete do Senador Eduardo Gomes e nos colocamos à disposição para outros diálogos com os Senadores sobre esses temas.

Muito obrigado.

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado, Sr. João Caldeira Brant Monteiro de Castro, Secretário de Políticas Digitais da Secretaria de Comunicação Social da Presidência da República.

Com isso, com a fala do Secretário, nós concluímos as falas de todos os apresentadores e agora, então, nós passamos para a fase dois aqui da nossa audiência, que começa com a leitura das perguntas que nós recebemos através do e-Cidadania. Essas perguntas foram enviadas pelo Portal do Senado, [www.senado.leg.br/ecidadania](http://www.senado.leg.br/ecidadania), e também pelo telefone 0800 0612211.

Então, eu começo já agradecendo também a participação de todos aqueles que nos acompanham pela TV Senado e daqueles que participaram aqui com as suas perguntas.

Eu passo à leitura das perguntas. Os debatedores que estão aqui já têm as perguntas em mãos. Para aqueles que estão remotamente, eu vou ler. Então, se vocês se interessarem em responder uma ou mais dessas perguntas, por favor, anotem, para a resposta.

Primeiro da Júlia Teresa, de Minas Gerais: "[...] [É possível citar] [...] casos em que a [...] [inteligência artificial] desempenhou um papel importante em eleições recentes?".

Novamente, Júlia Teresa, de Minas Gerais: "[...] [É possível citar] [...] casos em que a [...] [inteligência artificial] desempenhou um papel importante em eleições recentes?".

Pergunta dois, do Leonardo Cândido, de São Paulo: "Como garantir que o uso [...] [da inteligência artificial] em eleições seja transparente e ético, evitando [a] manipulação [...] [e] promovendo a inovação?".

De novo, Leonardo Cândido, de São Paulo: "Como garantir que o uso [...] [da inteligência artificial] em eleições seja transparente e ético, evitando [a] manipulação [...] [e] promovendo a inovação?".

Terceira pergunta, do Matheus Sales, do Rio Grande do Sul: "Uma vez que a inteligência artificial pode cometer equívocos, a exemplo dos atribuídos ao ChatGpt, como garantir a confiabilidade do uso desse sistema nas eleições?".

Novamente, Matheus Sales, do Rio Grande do Sul: "Uma vez que a inteligência artificial pode cometer equívocos, a exemplo dos atribuídos ao ChatGpt, como garantir a confiabilidade do uso desse sistema nas eleições?".

Quarta pergunta: Adriana Opstal, de São Paulo: "Como a inteligência artificial está sendo aplicada nas eleições e quais [...] as implicações dessa tecnologia na disseminação de informações?".

De novo, Adriana Opstal, de São Paulo: "Como a inteligência artificial está sendo aplicada nas eleições e quais são as implicações dessa tecnologia na disseminação de informações?"

Agora, então, conhecidas as perguntas enviadas pelo e-Cidadania, o que eu gostaria, dado nosso tempo.... Passamos o tempo, mas a discussão é boa, acaba passando. Isso é importante. É importante porque essas discussões nos dão aqui base para melhorar esse projeto de lei, mas também para que as pessoas, sociedade civil, participem e vejam tudo isso. É importante ter essa participação, afinal de contas aqui é a representação da nossa população.

Então, eu vou passar a palavra novamente para cada um dos nossos debatedores, para as suas considerações finais, por dois minutos cada um, e, ao longo dessas considerações finais, se quiserem responder uma ou mais dessas perguntas, fiquem à vontade, ou se quiserem colocar alguma pergunta para algum outro debatedor. *(Pausa.)*

O.k.

Então, eu vou começar aqui a nossa sequência de apresentações pelo Diogo Rais, que foi o primeiro, não é?

Diogo Rais, Diretor do Instituto Liberdade Digital, para as suas considerações finais e também para responder a algumas das perguntas.

**O SR. DIOGO RAIS** (Para expor.) - Perfeito, Senador.

Queria agradecer mais uma vez e dizer do prazer enorme em participar e aprender tanto com todos os colegas aqui.

Eu estava comentando com a Patricia da qualidade de todas as falas e do quanto nós temos vontade de aprofundar essas falas, indo além dos 10 minutos e além da voz dos 15 segundos, para que a gente possa aprender cada vez mais. Mas fica aqui toda a minha gratidão.

Eu só vou tentar aproveitar um instante. O Senador, logo após a minha fala, trouxe alguns dos dilemas para identificar a desinformação e como enfrentá-la sobretudo no grande dilema entre verdade e mentira. Como esse é um tema, Senador, com que eu tenho me ocupado há muitos anos também, de pesquisa, desde 2015, e tem me afligido muito a dificuldade em trazer uma conceituação *in abstracto*, eu passei a me dedicar ao olhar...

*(Soa a campanha.)*

**O SR. DIOGO RAIS** - ... de qual seria a função do direito em si nesse desafio. E, particularmente, tenho pensado que talvez desinformação, *fake news*, para que seja um problema jurídico ou como um problema jurídico, deveria ser considerada não uma falsidade em si, mas um conteúdo enganoso, o que é eventualmente mais amplo, e com potencial lesivo uma vez que o dilema puro entre a verdade e a mentira talvez ocupe mais pontos no campo da ética e da moral do que no direito em si. E me parece que o direito, de uma certa forma, serve para proteger os bens da vida. Então, talvez a ameaça de dano ou a existência de dano é o que deflagra o direito a entrar nisso e transforma a desinformação em um problema jurídico em si, e talvez esse comportamento nos ajude a enfrentar.

E, só para trazer um pouco de atenção aqui a todas as perguntas excelentes que foram enviadas, eu vou fazer um comentário muito breve a respeito da pergunta da Júlia Teresa, de Minas Gerais, também com gratidão por trazê-la aqui.

*(Soa a campanha.)*

**O SR. DIOGO RAIS** - Eu acredito que a inteligência artificial tem desempenhado um papel importante sobretudo no direcionamento de conteúdos digitais e, como eu disse antes, tanto para o bem quanto para o mal, também nas possibilidades, conexões e dificuldades que isso traz, como a própria colega Renata Mielli mencionou, a possibilidade de aumentar a polarização, intolerância e tudo mais.

E, por fim, porque eu já ultrapassei todas as minhas chances aqui, eu só gostaria de deixar um livro sobre inteligência artificial, porque, junto com a Profa. Juliana Abrusio, nós conduzimos uma disciplina de mestrado e doutorado na Universidade Mackenzie, na qual a gente abordou os temas da privacidade, mercado e cidadania. De uma certa forma, é um produto de toda a sala, e gostaria de deixar aqui com toda a gratidão por este momento, e me manter à disposição, e também me comprometer em enviar os materiais por escrito.

Muito obrigado, Senador.

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Eu que agradeço, agradeço pela participação, agradeço pelas respostas, pelos comentários e também agradeço pelo livro, vou deixá-lo na Comissão. E, com certeza, tudo isso vai ser muito importante para todos nós aqui. Obrigado.

**O SR. DIOGO RAIS** - Obrigado.

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Na sequência, então, para suas considerações finais e resposta a algumas das perguntas, eu passo a palavra para a Sra. Celina Bottino, Diretora de Projetos do Instituto de Tecnologia e Sociedade. Mantenha a palavra por dois minutos. Obrigado.

**A SRA. CELINA BOTTINO** (Para expor. *Por videoconferência.*) - Bom, então, fazendo coro aqui aos colegas, agradecendo muito este momento aqui, acho que não só a Comissão aprende, mas todos aqui estão ouvindo, não é? Podemos ter esta oportunidade de ouvir dos colegas todos os pontos de atenção, que não são poucos, não é? Mas, antes de talvez responder, tentar endereçar uma das perguntas, queria fazer um comentário.

Acho que a fala da colega Dora Kaufman, com a ideia dessa sugestão de criação desse conselho composto pelas agências reguladoras com os representantes, e, além disso, representantes de órgãos técnicos para discutir a eventual aplicação, para ver o que está sendo proposto - como que isso, vamos dizer assim, entraria; para que seria aquela aplicação na prática; quais seriam os desafios; e quais outros pontos que ainda precisam ser mapeados -, acho que esse exercício de trazer mais informação, mais dados, mais pesquisa empírica deve ser fortalecido.

(*Soa a campanha.*)

**A SRA. CELINA BOTTINO** (*Por videoconferência.*) - E, como a professora comentou, a discussão na Europa levou anos, e com muito mais informação do que a gente está tendo aqui, e, claro, estamos contra o tempo, mas, só reforçando essa ideia de trazer mais pesquisas e debates práticos sobre essa questão.

E, bom, respondendo, talvez, à pergunta do Leonardo, sobre como garantir que o uso da IA seja transparente e ético, garantir, garantir acho que ninguém vai dar 100% de certeza, mas, como a gente comentou, tem já medidas aí de marcas d'água, etc. e tal, mas não dá para confiar 100%, até porque a tecnologia avança no sentido de as pessoas colocarem também e até tirarem, eventualmente, depois essas marcas.

Então, isso tem um limite, como o próprio Diogo comentou, e aí, então, acho que nos resta reforçar todas as medidas de conscientização e análise crítica, porque isso acho que ninguém tira das pessoas.

(*Soa a campanha.*)

**A SRA. CELINA BOTTINO** (*Por videoconferência.*) - Então, termino por aqui, mais uma vez parabenizando esta audiência aqui pela Comissão.

Obrigada.

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado, Sra. Celina Bottino, Diretora de Projetos do Instituto de Tecnologia e Sociedade.

Eu passo agora a palavra, para as suas considerações finais, à Profa. Luciana Moherdaui, pesquisadora do grupo Jornalismo, Direito e Liberdade da Escola de Comunicações e Artes da Universidade de São Paulo, por dois minutos, e também para responder a algumas das perguntas que achar interessante.

Obrigado.

**A SRA. LUCIANA MOHERDAUI** (Para expor. *Por videoconferência.*) - Bom, eu agradeço, mais uma vez, o convite para essa audiência pública.

As falas dos colegas foram importantíssimas. Achei muito importante o que o João Brant disse sobre os projetos conversarem. Eu até incluí isso no *briefing*, mas o meu *briefing* ficou enorme, mas isso é um ponto crucial para que não haja repetição. Eu só não sei se questões ligadas à lei eleitoral têm que ficar necessariamente em projeto de lei eleitoral. Eu acho que o assunto, eu penso que esse assunto contempla tantas áreas que eu não sei se dá para ficar fatiando, enfim.

A respeito das perguntas, há, sim, iniciativas positivas, recomendo. O *The Washington Post* tem feito muitas reportagens sobre isso, apontando os pontos positivos do uso de IA em campanhas eleitorais, como eu falei aqui, principalmente em campanhas em que os candidatos não têm dinheiro para contratar grandes equipes de *marketing*. Enfim, esse é um ponto positivo.

Do ponto de vista do controle, como a Celina disse, mesmo com a marca d'água, é possível remixar o conteúdo e fazer algo novo, isso se faz sempre. A discussão sobre combate à desinformação vai por esse lado, porque, no momento em que uma informação passa pelo Twitter e uma conta é bloqueada, ela vai pulando - uma desinformação, perdão! -, ela pula de uma plataforma para outra, o que torna praticamente impossível conter tudo. Então, é complicado.

Agora, é preciso dizer, como jornalista, que eu acho importante que as pessoas apurem, que elas desconfiem das informações que recebem. Elas precisam procurar fontes confiáveis, enfim. Vou dar um exemplo agora: uma jornalista da

BBC e uma colega criaram aqui, por conta da guerra de Israel, uma checagem visual. Então, procurar fontes confiáveis, ir atrás de informações que sejam credíveis, na minha opinião, é uma saída, em vez de depender somente de legislação e de decisões também das plataformas ou, por exemplo, do TSE, enfim.

Eu encerro aqui, mais uma vez, agradecendo o convite. Obrigada, Senador. A audiência foi excelente. Eu saio daqui muito bem impressionada.

Bom dia a todos.

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado. Obrigado Profa. Luciana Moherdaui, Pesquisadora do Grupo Jornalismo, Direito e Liberdade, da Escola de Comunicação e Artes da Universidade de São Paulo.

Eu passo agora, então, a palavra à Sra. Patricia Blanco, Presidente do Instituto Palavra Aberta, para as suas considerações finais e respostas a perguntas.

**A SRA. PATRICIA BLANCO** (Para expor.) - Muito obrigada.

Bom, foi uma manhã riquíssima. Eu acho que, das audiências públicas de que eu participei sobre esse tema, essa foi realmente muito completa.

E, seguindo as falas, eu quero saudar a Secretaria da Comissão, que colocou as falas de uma forma tão encadeada que parecia que a gente tinha combinado antes - um completou o outro -, fazendo com que a gente conseguisse talvez percorrer todos os problemas, os dilemas que nós temos na regulação de algo tão complexo num curto espaço de tempo, de dez minutos, como a Profa. Dora colocou.

Eu queria aqui, primeiro, me colocar à disposição, o Instituto Palavra Aberta e todos - a Mariana, a Daniela e o Bruno - que têm trabalhado ativamente nessa questão da discussão dos conteúdos principalmente pedagógicos para levar educação midiática digital, letramento informacional...

*(Soa a campanha.)*

**A SRA. PATRICIA BLANCO** - ... letramento algorítmico para as escolas, formando professores. Então, queria colocar toda a nossa pesquisa à disposição do Senado e também da Secretaria aqui da Comissão.

E queria responder à pergunta do Matheus Sales, do Rio Grande do Sul, que coloca: "Uma vez que a inteligência artificial pode cometer equívocos, a exemplo dos atribuídos ao ChatGPT, como garantir uma confiabilidade [...] [do uso nas] eleições?". Eu colocaria aqui mais do que nas eleições: no uso geral. A gente precisa avaliar a resposta que o ChatGPT nos oferece a partir do pensamento crítico. Eu vou bater nessa tecla, porque eu acho que isso é fundamental. Não dá para a gente discutir inteligência artificial ou, como a Renata bem colocou, que não é inteligência, não é? - inteligência é um fator humano -, trazendo para essa tecnologia, sem discutir a questão do pensamento crítico, do letramento informacional, do letramento algorítmico e da necessidade de empoderar o cidadão para que ele possa fazer, de fato, uma diferenciação entre conteúdos.

Então, eu queria deixar, de novo, a minha contribuição e dizer da minha disponibilidade em continuar contribuindo com esse debate.

E só um último ponto: queria fazer um convite a todos os participantes desta bancada para a gente continuar essa conversa. Vamos fazer um seminário? Vamos propor um debate para que a gente possa ficar discutindo mais aprofundadamente esses pontos, que eu acho que são extremamente relevantes. Eu acho que o Palavra Aberta, o ILD, o CGI.br e tantos outros, o InternetLab, poderíamos nos unir e organizar esse encontro com a participação do Senador, claro, com certeza. Por favor. Obrigada.

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado, Sra. Patricia Blanco, Presidente do Instituto Palavra Aberta. Parabéns pela participação e pelos comentários também!

Na sequência, eu passo a palavra, para suas considerações finais e respostas a algumas das perguntas, ao Sr. Marcelo Antônio Rech, Presidente da Associação Nacional de Jornais.

Por favor, tem a palavra por dois minutos.

**O SR. MARCELO ANTÔNIO RECH** (Para expor. *Por videoconferência.*) - Muito obrigado, Senador.

Eu acredito que exista uma convergência no sentido de que é necessária, sim, uma regulação. A grande discussão é como deve se dar essa regulação e quando deve se dar essa regulação. É um risco que nós estabeleçamos uma regulação que já nasça defasada, e essa tentação de querer regar o detalhe tende a produzir legislações defasadas.

Há um ano, em novembro de 2022, nós não enxergávamos os desafios que nós temos visto hoje. Seguramente, daqui a um ano, haverá outros desafios que nós não teremos enxergado, pelo menos não com a nitidez, no momento em que eles passam a eclodir, a se desenvolver.

Então, a minha sugestão é que a gente faça uma discussão no sentido de uma regulação abrangente, que seja capaz de absorver com velocidade, seja pela própria regulação, seja por uma agência reguladora que tenha a agilidade de estabelecer princípios e regramentos novos à luz da velocidade que vai acontecer, à velocidade da luz, que nós estamos vendo no desenvolvimento da inteligência artificial, levando-se em conta, obviamente, as melhores práticas que estão sendo implantadas aí, em outros centros de pensamento e de discussão.

Seguramente, é claro, essa regulação deve levar em consideração os princípios elementares da Constituição brasileira: o respeito às leis, o respeito à liberdade de expressão, o respeito à propriedade intelectual e assim por diante, e se adaptar a esse novo mundo com uma capacidade de velocidade que nós, até hoje, não temos visto. O próprio marco civil da internet, que é uma excelente lei, já tem vários elementos, vários artigos já defasados, alguns de alta necessidade, como o próprio art. 19, que nós temos discutido.

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado, Sr. Marcelo Antônio Rech, Presidente da Associação Nacional de Jornais.

Eu passo, então, agora a palavra... Ela está ali? Está. Checando... À Sra. Renata Mielli, Coordenadora do Comitê Gestor da Internet para suas considerações finais e resposta a algumas das perguntas.

Obrigado.

**A SRA. RENATA MIELLI** (Para expor. *Por videoconferência.*) - Obrigada, Senador. Eu queria agradecer, mais uma vez, a possibilidade de participar desta audiência pública, que acho que foi riquíssima. Acho que todas as exposições, de fato, trouxeram elementos complementares sobre os impactos do uso da IA para as democracias, para o jornalismo.

Eu queria apenas aproveitar estes minutos finais aqui também reforçando o que foi colocado pela Profa. Dora e também pelo meu colega João Brant sobre a necessidade de a gente ter muita cautela na aprovação desse projeto de lei de regulação de IA. Toda regulação de tecnologia é bastante delicada porque as tecnologias se desenvolvem, surgem novos desafios, e a legislação pode ficar rapidamente obsoleta. Nós precisamos buscar quais são os princípios, mas também não devemos ficar em princípios abstratos, que não tenham impactos concretos no enfrentamento às externalidades negativas que a IA pode trazer para a sociedade.

E me somo ao João no que diz respeito ao que nós estamos falando, de um tema que é transversal a muitas matérias. Ele diz respeito à integridade da informação, ele diz respeito às questões relacionadas à saúde, à higidez do processo eleitoral. Portanto, o diálogo entre propostas legislativas que tratam desses temas é essencial.

Eu queria colocar também o Comitê Gestor da Internet no Brasil à disposição desse diálogo, desses espaços. No âmbito do CGI, a gente tem procurado discutir, de forma mais aprofundada, as questões relacionadas à inteligência artificial e também estamos, no âmbito do CGI, responsáveis pela criação e consolidação do Observatório Brasileiro de Inteligência Artificial, que eu acredito que pode também trazer aí ideias, aportes interessantes para a continuidade desse debate.

Conforme colocou aqui a Patricia, estamos aí à disposição e sempre com muito interesse de dialogar.

Muito obrigada mais uma vez.

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado, Renata Mielli, Coordenadora do Comitê Gestor da Internet. Obrigado.

Agora eu passo a palavra, então - agora sim -, para a Profa. Tainá Aguiar Junquilha, Professora de Direito, Inovação e Tecnologia do Instituto Brasileiro de Ensino, Desenvolvimento e Pesquisa.

**A SRA. TAINÁ AGUIAR JUNQUILHO** (Para expor.) - Então eu vou, fazendo as considerações finais e antes de agradecer a participação, responder aqui em bloco às perguntas, agradecendo também a participação de todos.

Alguns casos em que a inteligência artificial pode afetar a democracia a gente comentou aqui, mas, em especial, quando se produzem *deepfakes*, quando ela reforça a desinformação - e isso atrapalha a questão jornalística -, e quando ela reforça também a polarização, e a gente tem visto isso no mundo inteiro acontecer. Então, por isso, as minhas sugestões são nesse caminho que, de fato, foi falado aqui hoje: estímulo ao letramento digital, à educação, estabelecendo obrigações preventivas a esses desenvolvedores, às avaliações de impacto, regras, de fato, de governança que sejam fiscalizadas por um órgão "multiagencial", diríamos assim...

(Soa a campanha.)

**A SRA. TAINÁ AGUIAR JUNQUILHO** - ... e que traga a confiabilidade dos cidadãos nesses sistemas, ao mesmo tempo em que desenvolva a soberania do país.

E aí, finalizando, eu queria agradecer a oportunidade de ter falado aqui e aproveitar também essa oportunidade, dizendo duas coisas, fazendo um convite a este seminário proposto, que já existe: ele vai ser realizado dias 9 e 10 no IDP, e a gente vai fazer um lançamento, eu e a Profa. Laura Schertel vamos lançar o Leia (Laboratório de Estudos em Governança de Inteligência Artificial), propondo, justamente no IDP, estudos para auxiliar essa tarefa, que é tão árdua.

E também quero aproveitar a oportunidade para entregar a nota técnica elaborada em relação a esse Projeto de Lei 2.338 pela Coalizão Direitos na Rede em mãos aqui. Vão ficar também exemplares para auxiliar nessa tarefa que a gente sabe que é tão árdua.

*(Soa a campanha.)*

**A SRA. TAINÁ AGUIAR JUNQUILHO** - A gente fez aí comentários que buscam auxiliar nessa tarefa árdua aí da consecução desse projeto que é tão importante para o nosso país.

Muito obrigada.

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado à Profa. Tainá Aguiar Junquilha, Professora de Direito, Inovação e Tecnologia do Instituto Brasileiro de Ensino, Desenvolvimento e Pesquisa.

Eu passo a palavra agora, para as considerações finais, à Profa. Dora Kaufman, Professora da Pontifícia Universidade Católica de São Paulo, nossa PUC de São Paulo. Professora, para suas considerações finais.

**A SRA. DORA KAUFMAN** (Para expor. *Por videoconferência.*) - Muito obrigada, Senador.

Começando pelas perguntas que foram feitas, eu acho, pelas minhas pesquisas, pelas minhas investigações - não é o meu foco -, que ainda não teve uma situação, de fato, significativa que tenha tido impacto real nas eleições, nos resultados das eleições com o uso da inteligência artificial.

Sobre a tão falada - e foi citada hoje também aqui - Cambridge Analytica, eu recomendo a leitura deste livro... Não dá para ver. Dá para ver? Eu vou tirar... É o livro *Network propaganda*...

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Se colocar um pouquinho mais para trás, dá... Coloque um pouquinho mais para trás o livro na... *(Pausa.)*

**A SRA. DORA KAUFMAN** (*Por videoconferência.*) - É, não vai. Bom, é um livro que se chama *Network propaganda*. Um dos autores, o principal, é Yochai Benkler. Ele foi publicado em 2018 e analisa detalhadamente, com muita pertinência, o processo eleitoral de 2016 e conclui que não teve interferência - ou praticamente não teve - nem da Cambridge Analytica nem de outros instrumentos de inteligência artificial. Foi uma matriz de variáveis que levou ao resultado de 2016. E, aqui no Brasil, também foram feitos alguns estudos mostrando... Às vezes se confundem tecnologias digitais com o uso de inteligência artificial. Isso não quer dizer... Eu acho que todos nós estamos alertas, porque tudo indica... Ou pelo menos existem as condições, as condições são dadas para se usar nas próximas eleições a inteligência artificial e confundir a situação. O que vai acontecer, como será usado? Eu acho que nenhum de nós sabe, mas eu acho que cabe ao Tribunal Eleitoral estar atento e se preparar para gerir, para interferir quando o uso da inteligência artificial for feito no sentido de confundir a eleição.

E é isso.

Eu agradeço a oportunidade. Coloquei a minha proposta, porque eu estou muito preocupada com o processo. Eu acho que, volto a dizer, nós não estamos preparados enquanto país, não estamos maduros para estabelecer neste momento o marco regulatório, apesar de que, também volto a insistir, acho que o debate é urgente. A urgência do debate é uma coisa, e a não urgência de estabelecer o marco relatório. Acho que precisamos discutir amplamente. O Senador comentou sobre a participação na Anatel e outras agências. Volto a dizer, uma coisa é participar de uma consulta pública com dez minutos, outra coisa é sentar para trabalhar em conjunto, em meses ou o tempo que for, para produzir alguma colaboração de fato consistente que ajude, inclusive, o Senado a estabelecer esse marco relatório da inteligência artificial, que é fundamental.

Muito obrigada a todos.

E estou sempre à disposição.

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado, Profa. Dora Kaufman, Professora da Pontifícia Universidade Católica de São Paulo, a PUC de São Paulo.

Agora, para suas considerações finais, eu passo a palavra ao Secretário João Caldeira Brant Monteiro de Castro, Secretário de Políticas Digitais da Secretaria de Comunicação Social da Presidência da República. Para suas considerações, Secretário, por favor.

**O SR. JOÃO CALDEIRA BRANT MONTEIRO DE CASTRO** (Para expor. *Por videoconferência.*) - Obrigado, Senador.

Eu ressaltaria, na verdade, o elemento principal, acho, da minha fala sobre a regulação de IA: que esse debate precisa acontecer, nós precisamos maturar o processo em um tempo adequado. Acho também que não tem uma urgência grande, mas nós precisamos fazer com que ele seja um marco protetivo de direitos e de inovação ao mesmo tempo, mas que ele não deve fechar espaço para que regulações específicas respondam a questões setoriais. Eu acho que deve se manter o espaço e a legitimidade dos regulamentos específicos e acho que o desafio é como fazer isso ao mesmo tempo garantindo a proteção e promoção de direitos, fazendo com que os regulamentos específicos também não possam diminuir e enfraquecer o espírito que os Senadores propõem e trazem para o debate desse PL de IA.

Nesse sentido, nós nos colocamos à disposição de novo para dialogar tanto com o Relator quanto com o conjunto de Senadores da Comissão de Inteligência Artificial para buscar os melhores caminhos. E tenho certeza de que, com meus colegas de Governo que também estão envolvidos nesse tema - a Renata Mielli, que está aqui, é do MCTI; nós temos a Estela Aranha, que é do Ministério da Justiça e está inclusive nomeada para o grupo que auxiliará a ONU nesse debate; nós temos um conjunto de atores do MGI também, do Ministério de Gestão e Inovação -, podemos, inclusive, encaminhar soluções conjuntas para apreciação desta Comissão e para tomar os melhores caminhos.

Muito obrigado pela possibilidade e pelo convite.

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado, Sr. João Caldeira Brant Monteiro de Castro, Secretário de Políticas Digitais da Secretaria de Comunicação Social da Presidência da República.

Nós tivemos aqui, por cerca de duas horas e meia, esse debate bastante proveitoso.

Eu gostaria de aproveitar estes momentos finais para agradecer a cada um dos nossos debatedores que estiveram conosco aqui presencial e remotamente. Gostaria também de agradecer a cada uma das pessoas que nos acompanharam através da TV Senado e das redes sociais aqui do Senado, assim como àqueles que nos acompanharam presencialmente na sala de audiência nº 7 aqui do Senado. E também quero agradecer à nossa Mesa pelo trabalho de acompanhamento e por toda assistência técnica para que tudo isso acontecesse da forma mais fluida e mais eficiente.

Finalmente, não havendo mais nada a tratar, eu declaro encerrada a reunião, com votos de um ótimo dia para todos.

Obrigado.

*(Iniciada às 10 horas e 09 minutos, a reunião é encerrada às 12 horas e 43 minutos.)*