



SENADO FEDERAL
Secretaria-Geral da Mesa

ATA DA 14ª REUNIÃO DA COMISSÃO TEMPORÁRIA INTERNA SOBRE INTELIGÊNCIA ARTIFICIAL NO BRASIL DA 1ª SESSÃO LEGISLATIVA ORDINÁRIA DA 57ª LEGISLATURA, REALIZADA EM 01 DE NOVEMBRO DE 2023, QUARTA-FEIRA, NO SENADO FEDERAL, ANEXO II, ALA SENADOR ALEXANDRE COSTA, PLENÁRIO Nº 7.

Às onze horas e nove minutos do dia primeiro de novembro de dois mil e vinte e três, no Anexo II, Ala Senador Alexandre Costa, Plenário nº 7, sob a Presidência do Senador Eduardo Gomes, reúne-se a Comissão Temporária Interna sobre Inteligência Artificial no Brasil com a presença dos Senadores Carlos Viana, Efraim Filho, Nelsinho Trad, Fabiano Contarato, Astronauta Marcos Pontes, Rodrigo Cunha, Alessandro Vieira, Alan Rick, Angelo Coronel, Mara Gabrilli, Sérgio Petecão e Mecias de Jesus. Deixam de comparecer os Senadores Styvenson Valentim, Veneziano Vital do Rêgo, Weverton, Vanderlan Cardoso, Chico Rodrigues e Laércio Oliveira. Deixa de comparecer a Senadora Daniella Ribeiro, conforme REQ 631/2023-CDir. Havendo número regimental, declara-se aberta a reunião. A presidência submete à Comissão a dispensa da leitura e a aprovação da ata da reunião anterior, que é aprovada. Passa-se à Audiência Pública Interativa, atendendo ao Requerimento nº 3, de 2023-CTIA, de autoria Senador Astronauta Marcos Pontes (PL/SP), com a finalidade de debater os impactos da Inteligência Artificial nos Setores da Indústria, Agricultura, Público, Financeiro e Judiciário, com a participação de Giovanni Cerri, Presidente do Conselho de Inovação do Hospital das Clínicas da Faculdade de Medicina da Universidade de São Paulo (FMUSP); Daniel Stivelberg, Coordenador do Grupo de Governança e Regulação de Dados da Zetta; Marco Antonio Lauria, Membro do Conselho da Associação Internacional de Inteligência Artificial (A2IA); Crisleine Barboza Yamaji, Professora do Instituto Brasileiro de Mercado de Capitais (IBMEC); Robert Janssen, Presidente da Associação das Empresas Brasileiras de Tecnologia da Informação do Estado do Rio de Janeiro (Assespro-RJ); Bruno Jorge Soares, Gerente da Unidade de Difusão de Tecnologias da Agência Brasileira de Desenvolvimento Industrial (ABDI); Walter Marinho, Coordenador de Governança em Ciência, Tecnologia e Inovação da Rede Governança Brasil (RGB); e Luiz Fernando Bandeira de Mello Filho, Conselheiro Nacional de Justiça (CNJ). Nada mais havendo a tratar, encerra-se a reunião às doze horas e cinquenta e seis minutos. Após aprovação, a presente Ata será assinada pelo Senhor Presidente e publicada no Diário do Senado Federal.

Senador Eduardo Gomes
Presidente Eventual da Comissão Temporária Interna sobre Inteligência Artificial no Brasil

Esta reunião está disponível em áudio e vídeo no link abaixo:

<https://www12.senado.leg.br/multimidia/evento/117991>

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO. Fala da Presidência.)
– Declaro aberta a 14ª Reunião da Comissão Temporária Interna sobre Inteligência Artificial no Brasil, criada pelo Requerimento nº 722, de 2023, com a finalidade de, no prazo de até 120 dias, examinar os projetos concernentes ao relatório final aprovado pela Comissão de Juristas responsável por subsidiar



SENADO FEDERAL
Secretaria-Geral da Mesa

a elaboração do substitutivo sobre inteligência artificial no Brasil – criada pelo Ato do Presidente do Senado Federal nº 4, de 2022 –, bem como eventuais novos projetos que disciplinarem a matéria.

Aprovação da ata.

Antes de iniciarmos os nossos trabalhos, proponho a dispensa da leitura e a aprovação da ata da reunião anterior.

As Sras. Senadoras e os Srs. Senadores que estiverem de acordo permaneçam como se encontram. (*Pausa.*)

Aprovada.

A ata está aprovada e será publicada no *Diário Oficial do Senado Federal*.

Pauta: informo que esta reunião se destina à realização de audiência pública com o objetivo de debater os impactos da inteligência artificial nos setores de indústria, agricultura, público, financeiro e judiciário, em cumprimento a Requerimento nº 3, de 2023, da Comissão Temporária de Inteligência Artificial, de autoria do Senador Astronauta Marcos Pontes.

O público interessado em participar desta audiência pública poderá enviar perguntas e comentários pelo endereço www.senado.leg.br/ecidadania ou ligar para 0800.0612211.

Encontram-se presentes no Plenário da Comissão o Dr. Luiz Fernando Bandeira de Mello, Conselheiro Nacional do CNJ, a quem convido para fazer parte da mesa, um grande amigo, com relevantes serviços prestados ao Congresso Nacional e ao Brasil; Giovanni Cerri, Presidente do Conselho de Inovação do Hospital das Clínicas da Faculdade de Medicina da Universidade de São Paulo; Bruno Jorge Soares, Gerente da Unidade de Difusão de Tecnologias da Agência Brasileira de Desenvolvimento Industrial (ABDI); Daniel Stivelberg, Coordenador do Grupo de Governança e Regulação de Dados da Zetta; Walter Marinho, Coordenador de Governança em Ciência, Tecnologia e Inovação da Rede Governança Brasil.

Encontram-se também presentes por meio de sistema de videoconferência: Marco Antonio Lauria, membro do Conselho da Associação Internacional de Inteligência Artificial, A2IA; Crisleine Barboza Yamaji, Professora do Instituto Brasileiro do Mercado de Capitais (IBmec); Robert Janssen, Presidente da Associação das Empresas Brasileiras de Tecnologia da Informação do Estado do Rio de Janeiro (Assespro-RJ).

Agradeço a presença a todos.

Faço, agora, as considerações iniciais sobre a audiência de hoje.

O nosso querido Presidente, Senador Marcos Pontes, encontra-se numa missão na Aeronáutica brasileira. Deve participar também por videoconferência – está tentando fazer isso.

Então, agradeço a ele o convite.

Como Relator, vamos acompanhar aqui as exposições.

Iniciando as exposições, por até dez minutos, para suas considerações, como primeiro inscrito, o Dr. Luiz Fernando Bandeira de Mello Filho, Conselheiro Nacional de Justiça (CNJ).

O SR. LUIZ FERNANDO BANDEIRA DE MELLO FILHO (Para expor.) – Muito bom dia a todas e todos.

Eu cumprimento, inicialmente, meu querido amigo Senador Eduardo Gomes. É uma satisfação, Senador, voltar a esta Casa.

Eu sou servidor da Casa e por aqui estive por 19 anos já. Estou, por enquanto, licenciado, no CNJ, por delegação de V. Exas. Mas é um prazer muito particular, sendo servidor da Casa, retornar aqui para ajudar a instruir um projeto de lei. Tive a oportunidade recentemente de estar na CCJ, falando do



SENADO FEDERAL
Secretaria-Geral da Mesa

projeto de lei do *impeachment*, cuja experiência vivenciamos juntos; e, agora, sobre inteligência artificial.

O que me traz aqui hoje, Senador, e acabo...

Perdão, antes de iniciar, cumprimento também meus colegas de mesa e todos os servidores, meus colegas do Senado Federal, que cumprimento na pessoa da Adriana Nunes – a todos vocês. É um prazer revê-los.

O que me traz aqui hoje, Senador, é a condição particular de, enquanto Conselheiro Nacional de Justiça, eu ser responsável, no âmbito do CNJ, pela Comissão de Tecnologia da Informação.

É um tema que me é muito caro, até porque, antes de fazer meus estudos em Direito, eu era programador de computadores. Então, sempre tive muito interesse na área, e, chegando ao CNJ, foi-me delegada essa missão, que vem justamente a calhar e casar-se com a temática deste PL, que é a regulação da inteligência artificial.

Vou fazer algumas considerações iniciais sobre o pé em que estamos no CNJ e no Poder Judiciário como um todo acerca da discussão da inteligência artificial, para depois fazer considerações em cima do PL, sempre respeitando o tempo originalmente proposto.

Inteligência artificial não é novidade, já sabemos todos nós. Ela já vem sendo aplicada mundo afora em diversas aplicações, desde previsão do tempo até controle de semáforos nas cidades. A grande novidade, o que realmente chamou a atenção e assustou um pouco a todos nós, foi, a partir de novembro do ano passado, quando se tornou disponível uma ferramenta de inteligência artificial generativa, que é essa que produz textos ou imagens; mas, em particular, o ChatGPT, que tornou essa experiência acessível ao usuário comum, ao usuário leigo. Isso nos surpreende todos, porque, de repente, nós vemos uma máquina capaz de produzir um texto em menos tempo e talvez com melhor qualidade do que nós mesmos faríamos. Rapidamente pensamos: "Meu Deus, vou ficar obsoleto?", e a gente pensa logo nas consequências, inclusive profissionais, disso.

A partir do surgimento do ChatGPT, uma série de especulações surgiram na sociedade em geral. No que se refere ao Judiciário, que é particularmente o que eu vou tocar aqui, o grande medo é: vão usar o ChatGPT para julgar os casos que estão em andamento no Poder Judiciário? Inclusive, um cidadão provocou formalmente o CNJ a fim de proibir a utilização do ChatGPT no âmbito do Judiciário. Por acaso, dei parecer nesse caso, já é público, por isso posso aqui comentar.

Porém, a tônica que eu quero abordar com as senhoras e com os senhores não é propriamente a utilização do ChatGPT no Poder Judiciário, e sim da inteligência artificial generativa. Isso porque o ChatGPT é uma aplicação, é um programa colocado à disposição – tem vários outros concorrentes inclusive, produzidos pela Google, por exemplo, pela Meta, pela própria Amazon, que também vem desenvolvendo.

O grande diferencial de uma ferramenta de IA generativa aplicada a um setor como o Judiciário é a base de dados que a alimenta. O ChatGPT foi montado com base nos textos disponíveis livremente na internet – estou falando de blogues, estou falando de *sites* jornalísticos, estou falando eventualmente de redes sociais. Com base naquilo ali, ele se alimenta evidentemente com uma fonte de língua inglesa muito maior do que a que temos em português disponível. E ele consegue, através de um exame probabilístico, matemático quase, montar os *tokens*, que são, praticamente, sílabas que, conforme a probabilidade de estarem juntos em um determinado tema, fazem sentido para o leitor ordinário. Mas aquilo ali não é uma máquina pensando, é a máquina apenas aplicando um padrão de linguagem que nós já vemos repetido em outras estruturas.



SENADO FEDERAL
Secretaria-Geral da Mesa

O problema que várias pessoas apontaram é: "O ChatGPT foi usado nos Estados Unidos e inventou precedentes jurisprudenciais falsos...", quando a máquina alucina ou delira, criando, na verdade, em cima daquela experiência que ela tem representada na sua base de dados.

A ideia, quando estamos discutindo a inteligência artificial generativa no Poder Judiciário é, evidentemente, trabalhar com uma ferramenta que seja treinada com a base de dados adequada, para começo de conversa. Não dá para você pensar em usar uma ferramenta genérica como essas que estão disponíveis gratuitamente, até porque elas servem para bater papo. Na verdade, o ChatGPT é isto, um *chat*, um diálogo quase com a máquina, que tem, sim, uma capacidade impressionante de produção de texto, mas não é treinada para dar uma resposta jurídica.

Essa base de dados nós já temos. Isso talvez seja uma coisa que pouca gente saiba. O CNJ já vem desenvolvendo, há dois anos, uma base de dados gigantesca chamada Codex. Na verdade, o Codex é um robzinho que alimenta um *Data Lake* – o pessoal dessa área adora falar inglês, não é? É um lago de dados, digamos assim, em uma tradução literal – e esse robô vai entrar em cada processo judicial brasileiro que esteja no PJe, o sistema nacional de gestão eletrônica dos processos, e vai ler cada petição e cada decisão ocorrida nesse processo. Não estou falando do andamento processual, não, o andamento tem milhares de ferramentas que já checam isso e geram inclusive os *push* para os advogados acompanharem o processo. Estou falando do conteúdo das peças e das decisões judiciais. Então, o Codex lê o conteúdo de todos os processos e reúne isso em uma base de dados gigantesca. Isso já existe, isso já foi feito.

Quando eu falo em aplicar a inteligência artificial para treinar uma máquina dessa, eu não estou falando em 300 acórdãos ou decisões judiciais, estou falando de 30 milhões, a quantidade e a proporção são muito maiores, e nós temos uma base de dados já desse tamanho.

O que nós não temos? Nós não temos ainda o *software* treinado para usar essa base de dados e gerar texto com a qualidade e com a perspectiva jurisprudencial nacional, mas nós vamos ter.

O Presidente Luís Roberto Barroso, no dia da sua primeira sessão no CNJ, falou, publicamente, e eu posso, então, reproduzir aqui, que ele já manteve contatos com as *big techs* Meta, Google, Microsoft e Amazon, para que desenvolvam pilotos para o Poder Judiciário que permitam classificar os processos. Já temos isso, hoje, em funcionamento no STJ, no STF, em vários tribunais do país, robôs que usam inteligência artificial para classificar aquela demanda quanto ao tema e quanto a espécies de pedido, para que se possa colocar na caixinha correta, para o magistrado apreciar, mas, além de classificar, passar a propor soluções para o caso, com base na jurisprudência nacional. Não é uma decisão aleatória do computador. Ninguém pensa em permitir que robôs julguem os casos colocados ao escrutínio do Poder Judiciário. A ideia é facilitar o trabalho de redação de minutas, porque, fatalmente, na hora em que você vai julgar um caso, terá uma análise da prova, do caso concreto, mas existirá o momento de redação propriamente de uma minuta a ser finalmente assinada pelo magistrado. Essa minuta pode ser redigida pelo próprio magistrado ou, eventualmente, por uma assessoria. Em geral, nos tribunais, o que se tem é uma assessoria qualificada que desenvolve a minuta no estilo do desembargador e do ministro. Isso é assim em todo lugar do mundo. Não estamos inventando a roda. Mas qual a diferença? É que, quando eu vou fazer uma decisão, por exemplo, sobre licitações e contratos administrativos, eu quero citar um determinado doutrinador, como Maria Sílvia Di Pietro, Celso Antônio Bandeira de Mello...

(*Soa a campainha.*)

O SR. LUIZ FERNANDO BANDEIRA DE MELLO FILHO – Um minutinho. Vou tentar abreviar.



SENADO FEDERAL
Secretaria-Geral da Mesa

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO) – O senhor tem o tempo de que precisar. Esse dispositivo foi o senhor que criou aqui. *(Risos.)*

O SR. LUIZ FERNANDO BANDEIRA DE MELLO FILHO – É verdade. Estou sendo vítima da própria eficiência do sistema.

Naturalmente, você vai buscar citar alguma referência dogmática ou jurisprudencial, e isso leva um tempo, um tempo em homem e hora de trabalho, porque eu vou ter que buscar o livro na minha prateleira, ou, se eu o tiver em meio eletrônico, vou ter que fazer uma busca dentro do texto dele para buscar aquele trecho que eu desejo, eventualmente transcrever ou copiar e colar na decisão, eventualmente, o magistrado, quando for revisar aquela decisão, pode dizer "não, eu não queria citar essa passagem aqui da Maria Sylvia, eu queria citar aquela outra do professor fulano..."

(Interrupção do som.)

O SR. LUIZ FERNANDO BANDEIRA DE MELLO FILHO – Isso aí fui eu que criei mesmo. *(Risos.)*

Que desastre! Veja o perigo da máquina quando a gente a deixa agir sozinha.

E você, naturalmente, perderá um tempo para desenvolver aquela minuta.

Qual é a ideia do sistema que se pretende para o Judiciário? É que, uma vez parametrizado, ele produza um texto no estilo do magistrado, até porque ele vai usar como base as próprias decisões do magistrado em processos anteriores, e traga lá uma citação, por exemplo, do Ministro Celso de Mello.

E aí, você, quando vai revisando aquilo, pensa "não, Celso de Mello talvez seja um pouco antigo, eu queria uma jurisprudência de um ministro atual da Corte Suprema". Vai ter um botãozinho lá para gerar um novo, e ele vai buscar na base de dados outra citação compatível com aquele texto, e o magistrado vai decidir se quer usar essa ou, então, ainda uma outra.

Então, a ideia essencial é que a inteligência artificial sirva para recolher informações, porque, para isso, ela é muito melhor que o ser humano, ela é muito mais rápida e mais capaz do que o ser humano, mas, fatalmente, necessariamente, ela precisará submeter-se ao escrutínio do olhar humano para, só a partir daí, virar uma decisão judicial.

Com isso, a gente tem a convicção de ganhar um tempo enorme nessa análise de processos e, evidentemente, conseguir parametrizar melhor decisões, para evitar decisões conflitantes, para evitar raciocínios diversos. É claro que o magistrado ainda pode pensar diferente da média da jurisprudência nacional, ele ainda poderá divergir. A IA poderá sugerir para ele uma decisão, mas ele pode também instruir a IA para produzir uma minuta em outro sentido.

Hoje, a Advocacia-Geral da União já vem utilizando inteligência artificial no Sapiens. Eles usaram a versão do ChatGPT aplicada à base de dados dele. O sistema, por trás do usuário, manda instruções, o contexto, como é dito na área, o contexto do que se deseja – "Olha, essa é uma peça de representação judicial da União, o que se deseja é contestar a suposta prescrição do crédito tributário", algo assim –, vai passar instruções para a máquina, e a máquina, com base naquelas instruções, produz uma minuta, que o advogado da União irá ler, revisar e, eventualmente, aprovando-a, dar ingresso no Judiciário. Isso já existe do lado da AGU, é necessário que exista do lado do Poder Judiciário.

Eu não vou me estender demais, porque eu quero... é importante ouvir os demais colegas; mas a ideia fundamental não é, portanto, ter um robô julgando os casos, é ter um robô ajudando o juiz a produzir uma decisão, que ele já teria, racionalmente falando, mas, em vez de digitar talvez umas seis, oito, dez páginas de decisão, ele consegue ter um levantamento de informações com base na jurisprudência do próprio magistrado e de outros tribunais superiores e produzir um texto mais facilmente, mais rapidamente.



SENADO FEDERAL
Secretaria-Geral da Mesa

Não é possível ficar para trás nessa matéria. Inteligência artificial é o futuro em diversos setores da economia e do Estado, e não faz sentido abrimos mão disso, desse ganho tecnológico.

Em relação ao PL, Sr. Senador, eu tenho certeza de que poderemos falar um pouco mais dele adiante, mas o mais importante é não fechar portas. Nós temos, por exemplo, no CNJ, uma regulamentação sobre inteligência artificial de três anos atrás e que já precisará ser toda revista. E digo por quê: porque nela se previa que o Judiciário seria quase que uma *software house*. Nós íamos desenvolver a nossa inteligência artificial, o que é uma coisa caríssima e demoradíssima e que possivelmente não é da natureza, a atividade-fim do Poder Judiciário, que é julgar, e não produzir *software*. Então, nós teremos que atualizar essa norma do CNJ, a fim de permitir que se contrate, seja com uma *big tech* dessas internacionais, seja com alguma empresa nacional de desenvolvimento de IA – a UnB vem fazendo um projeto importante nessa área –, e certamente iremos contratar algum serviço desses para, aplicado sobre a base do Codex nacional, você ter um sistema que realmente atenda o Poder Judiciário brasileiro. Nossa expectativa é que, em seis meses, tenhamos um piloto funcionando, para que, quem sabe, num prazo de um ano e meio, talvez dois anos, isso esteja já operacional para todo o Brasil. Envolve custos, envolve orçamento – porque evidentemente isso tudo se paga –, mas é um desafio que nós temos que enfrentar, porque senão vamos ficar para trás. E a ideia, sem dúvida, é avançar, e acho que o projeto funciona muito bem para isso.

Finalmente, como última consideração ainda acerca do projeto, é importante deixar claro que, quando ele for usado para servir ao Estado, seja no Judiciário, seja na AGU, seja eventualmente na Secretaria do Tesouro Nacional, sempre exista a responsabilidade do operador, da autoridade pública que está usando aquele sistema, que, ao adotá-lo, seja responsável pelo que sai dali – até mesmo para obrigar as autoridades a revisarem efetivamente o texto. Ninguém quer deixar as máquinas decidirem os rumos do Estado brasileiro.

São algumas considerações iniciais, Senador. Eu ficarei aqui à disposição para voltarmos à temática. Obrigado pela oportunidade.

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO) – Muito obrigado ao Dr. Luiz Fernando Bandeira de Mello Filho pela exposição.

Eu faço rapidamente aqui um registro importante para ficar à disposição dos nossos expositores, para interação e resposta de perguntas que foram feitas pelo sistema e-Cidadania.

O Jefferson Azevedo, de São Paulo, pergunta: "A invasão de privacidade é uma questão importante. Os legisladores têm o papel crucial de equilibrar o avanço da tecnologia da inteligência artificial?".

O Carlos Henrique, do Maranhão, pergunta: "Existe projeto para desenvolver a educação continuada dos trabalhadores [...]?".

O Everton Cassiano, do Rio Grande do Sul: "Como a IA pode otimizar produtividade e ética, garantindo inovação [...] nos setores industrial, agrícola, público, financeiro e jurídico?".

Tatiane Pereira, do Rio de Janeiro, pergunta: "A tecnologia facilita a produção em 'massa', mas como identificar [as áreas em que] a IA não deve invadir? [...]".

Vinícius Bergamoni, de Minas Gerais, faz um comentário aqui: "O futuro exige que as IAs automatizem processos judiciais, auxiliem nos cálculos financeiros [...]".

Marcio Nonato, do Pará: "Com toda certeza terá impacto positivo em todos os setores, mas [...] [melhor seria] seria começar pelos [...] [setores] de pesquisa e universidades".



SENADO FEDERAL
Secretaria-Geral da Mesa

Thaynara Duarte, de Minas Gerais, também deixa seu comentário: "No contexto acadêmico, é crucial impulsionar pesquisas inovadoras e preparando estudantes para os desafios e oportunidades futuras".

Apenas para passar para o próximo expositor, fazer o comentário de que esta Comissão Especial conclui no dia de hoje 80 expositores, que fizeram sua participação através das audiências públicas, durante todos esses dias, e, por ordem do nosso Presidente Rodrigo Pacheco, teremos ainda a realização daquilo que na Câmara se chama comissão geral, aqui é sessão de debates, para preparar a ida do PL a Plenário e votação nas Comissões, entendendo que já foi cunhado aqui o termo, única coisa que a gente encontra para explicar, que teremos que fazer uma lei viva, minimalista, mas eficiente. Tivemos o exemplo recente, por exemplo, da lei de proteção de dados. V. Exa. nos ajudou e conseguimos aprovar a PEC 115.

Passo a palavra, neste momento, ao Bruno Jorge Soares, Gerente da Unidade de Difusão de Tecnologias da Agência Brasileira de Desenvolvimento Industrial (ABDI).

O SR. BRUNO JORGE SOARES (Para expor.) – Bom dia a todos.

Primeiramente, eu gostaria de agradecer o convite do Senado Federal na pessoa do Senador Eduardo Gomes e cumprimentar todos que estão presentes aqui. Para mim, é uma satisfação e uma honra muito grande contribuir com este debate em relação ao futuro do Brasil, principalmente quando a gente pensa numa tecnologia tão importante como a IA.

A ABDI (Agência Brasileira de Desenvolvimento Industrial) é uma agência vinculada ao Ministério do Desenvolvimento, Indústria e Comércio. A gente tem como foco o aumento da produtividade e avanço da indústria no Brasil, seja indústria no sentido lato, no setor econômico, no agro, na saúde e nos serviços, em todos os segmentos em que a atividade econômica possa contribuir para a sociedade.

Nesse sentido, eu queria destacar na minha fala que a indústria no mundo hoje corre numa jornada, num avanço de um binômio entre a tecnologia e a inovação, e que a inteligência artificial é basicamente a linha de chegada de uma nova etapa dessa indústria. A gente pode dividir a indústria em gerações, a gente fala da primeira revolução industrial, da segunda revolução industrial, e a quarta revolução industrial, indústria 4.0, que é o tema que a gente trabalha dentro da ABDI e é onde a inteligência artificial alcança sua plenitude.

Mas, Senador, essa corrida não é fácil para as indústrias, não é fácil para o Governo, no sentido de que o processo de desindustrialização que a gente viveu e de transformação social e econômica do setor industrial colocou essa questão da tecnologia como um fator preponderante para o sucesso e a sobrevivência da indústria.

A indústria no Brasil, se a gente quiser falar no futuro de 2030 e muitos anos ainda pela frente, vai depender da tecnologia para a sobreviver. E basicamente a tecnologia que vai mudar esse padrão é a inteligência artificial. Em alguns levantamentos que a Agência Brasileira conduziu junto com seus parceiros, fornecedores, a gente comprova que o uso ainda da inteligência artificial acontece de forma isolada nas empresas no Brasil.

Eu falo do ponto de vista da manufatura, basicamente, que é o tema que a gente mais trabalha. Quando a gente pensa no ciclo da inteligência artificial, o Luiz Fernando, do CNJ, colocou a questão dos dados aqui como um passo fundamental.

É como se as empresas no Brasil, de uma maneira geral, na média, considerando as micro, pequenas e até uma parte das médias e grandes, não avançassem na maturidade desses dados. É como se a gente precisasse fazer um dever de casa muito grande no Brasil como um todo.



SENADO FEDERAL
Secretaria-Geral da Mesa

Claro que existem exceções e vários exemplos no Brasil de aplicações de IA e esses exemplos é que vão para a mídia e chegam ao conhecimento de V. Exas., os Senadores e outros Parlamentares, como no sentido do que está acontecendo na IA, mas isso não é uma verdade para a indústria como um todo. Pelos nossos levantamentos, a gente tem o indicativo de que apenas 2% das indústrias brasileiras estão na quarta geração da indústria 4.0.

Então, o que a gente está falando é discutindo uma regulamentação que precisa de um pé, além de ser uma lei viva, ela precisa de um braço de fomento muito importante. A gente não pode pensar essa questão da tecnologia como um todo, da IA, avançando sem o braço do fomento, sem o braço do incentivo, sem o braço da pesquisa e desenvolvimento, como foi destacado pelas perguntas feitas pela plateia *online*. E é neste sentido que a ABDI tem trabalhado: como fomentar projetos nessa área.

Então, em um projeto recente sobre indústria para planejamento da produção, o uso de IA, num edital que teria recursos não reembolsáveis, que nem alguns dizem: "a fundo perdido", a empresa não precisa reembolsar a agência ou o ministério ou qualquer ente público pelo uso dos recursos – ela apresenta um projeto e esse projeto é beneficiado com um valor –, nós tivemos em 2020 esse edital de IA na produção e tivemos apenas 13, 14 empresas candidatas a esse projeto.

Claro que tem outros requisitos, mas isso mostra que, mesmo com incentivo, você tem ainda a questão da maturidade, de um lado, das empresas para aplicar e um conhecimento que ainda precisa ser difundido sobre a IA. A gente está em Brasília, estamos nos grandes centros, esse tema parece que é pervasivo à sociedade, mas muitas empresas desconhecem como a IA pode mudar seus negócios. Então, do lado do fomento, eu teria uma parte de difusão e educação da sociedade e do setor produtivo quanto aos benefícios da inteligência artificial.

Esses dois aspectos, para mim, são muito importantes quando a gente tem inspirações de regulação baseadas, por exemplo, na União Europeia, que tem uma realidade de empresas nesse patamar muito diferente do Brasil. É como se a gente estivesse trazendo regras de um jogo, de um campeonato que a gente ainda não joga na plenitude. E isso preocupa. Aqui falando pelo lado das empresas, com que a gente tem contato direto, essa discussão da regulamentação sem um braço do fomento e sem um braço da difusão.

A difusão do conhecimento da IA, dos benefícios da IA, eu acho que em vários setores, a inteligência artificial já avançou bastante. A saúde é um setor de que o Prof. Giovanni, com quem a gente tem projeto lá, com o HC da USP, vai poder falar também.

No agro, temos também avanços; na indústria, temos avanços. Mas quando a gente pega o todo da indústria, isso precisa ser discutido.

E eu acredito que o principal recado é uma lei viva, que fomenta e que educa a sociedade. Essa é a minha mensagem. Eu acho que é uma expressão importante, uma orientação. E coloco a ABDI à disposição para reportar esses estudos, esses dados, à disposição desta Comissão e do Senado Federal para contribuir com esse debate.

Obrigado.

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO) – Obrigado, Bruno Soares, pela exposição.

Passo a palavra agora ao Dr. Walter Marinho, Coordenador de Governança e Ciência, Tecnologia e Inovação da Rede Governança Brasil (RGB).

Com V.Sa. a palavra.



SENADO FEDERAL
Secretaria-Geral da Mesa

O SR. WALTER MARINHO (Para expor.) – Muito bom dia a todos. Agradeço imensamente ao nosso Relator e ao Senador Eduardo Gomes por nos receber. E especialmente ao Senador Marcos Pontes pelo convite. Estou aqui a convite dele.

E vou fazer algumas pontuações e reflexões sobre governança e ética em inteligência artificial.

Eu morei por muitos anos na Europa, trabalhava com políticas estratégicas entre a Europa e os países de língua portuguesa, e até algumas perguntas que cá vieram vão fazer reflexiva no que eu aqui descrevi para nós podermos defender.

Mas já se sabe que o conhecimento é poder desde sempre, não é? E antigamente as questões morais eram um dos critérios decisórios para ter acesso a esse conhecimento, quando o indivíduo necessitava de formação antes de ter acesso a tal poder. Isso era justificável, porque era a fim de mitigar o risco de uma revolução tecnológica sobrepor a revolução humana, ou seja, o indivíduo não estar preparado.

E o que isso tem a ver com nós estarmos aqui hoje? Tem tudo a ver, devido às nossas revoluções tecnológicas, em que hoje um indivíduo qualquer, em qualquer parte do mundo, por um clique ou por um sistema de voz, tem acesso à informação para fazer qualquer coisa, de qualquer forma. Fazendo uma certa analogia, é uma criança com uma arma nas mãos.

Então, nós temos um amplo uso da tecnologia, à qual a governança e a ética devem estar incrustadas para que nós possamos mitigar riscos e vieses dessas tecnologias.

Antes de que nós possamos avançar, o PL trouxe grandes contribuições que, hoje, nos fazem estar aqui. Mais de 80 expositores já cá estiveram, como se mencionou. Devemos entender o que seja a inteligência artificial. A forma como ela consta lá, eu acredito que não esteja na minha visão correta. A inteligência artificial nada mais é do que um modelo matemático, estatístico, aplicado, evolutivo, e devemos ter um olhar mais atento e profundo sobre a sua aplicação em governos e empresas e às causas que nos fazem tomar decisões nas diversas áreas de sua aplicação. É muito importante nós termos esse olhar atento e termos uma certa visão justificativa do *design* que nos faz tomar essas decisões. É muito mais importante ainda nós tirarmos nossas limitações de acharmos que a máquina imita as funções cognitivas dos seres inteligentes, que raciocina, que aprende e resolve problemas. É muito importante para o nosso debate.

E isso tem consequências. Tanto é que 155 países se reuniram em um documento para estabelecer o que seja ética e governança de inteligência artificial durante a pandemia. O resultado disso foi que eles temiam que o desenvolvimento da IA pudesse nos tirar a capacidade humana de decidir durante o combate à pandemia. Lembrando que a inteligência artificial é um sistema computacional limitado que não pode ir além do que foi programado.

O preconceito pode ser colocado em qualquer fase de qualquer projeto de inteligência artificial. Então, os dados não são neutros, como se pensa. Assim como a gente define, em algum momento, que a gente decide usar um dado ou não, nós já podemos nos aproximar de algum viés. E o preconceito não é necessariamente intencional nem imediatamente evidente, mas a governança e a ética podem detectar e ajudar a mitigar esses riscos e eliminá-los desde o início do projeto.

A inteligência artificial é incontornável e os seus desdobramentos sobre isso são muito importantes. A Inglaterra utilizou o sistema de inteligência artificial para poder escolher, durante a pandemia, quem seria atendido primeiro. Com isso, nesses casos, muitas pessoas que tiveram este atendimento postergado vieram a óbito. Qual a consequência desse ato? É que os familiares que foram afetados por essa inteligência artificial não sabiam quais critérios foram utilizados para que aquele



SENADO FEDERAL
Secretaria-Geral da Mesa

paciente fosse atendido daquela forma ou não. Então, isso é uma consequência em que a transparência, a governança e a ética deveriam estar mais presentes para dar tranquilidade a essas pessoas.

Os governos hoje funcionam integrando a sua atividade, conjuntamente, com as empresas privadas. Como o Governo pode assegurar aos cidadãos a ética e a governança desses serviços de atendimento? Garantindo que os dados do cidadão não estejam sendo usados para atos escusos, para fins políticos, religiosos ou mesmo até em situações que possam constranger a identidade daquele cidadão do ponto de vista da justiça ou da coerção.

Hoje o cidadão não tem visibilidade sobre o contexto de como são tomadas essas decisões. Então, a governança e a ética podem ajudar nessa transparência.

E, falando um pouquinho sobre o que o Sr. Daniel mencionou, temos alguns casos no mundo sobre o uso da inteligência artificial no âmbito do Judiciário, em que os tribunais possuem hoje sérias acusações de erros éticos resultantes de tomadas de decisões automatizadas de inteligências artificiais. Essas acusações vão de falsas dívidas a supostamente sentenças de prisões com preconceitos raciais, religiosos, e até políticos em algumas ditaduras digitais que nós temos aí pelo mundo.

Então, a governança e a ética deveriam dar diretrizes a esses projetos para evitar tais preconceitos e vieses na geração de transparência nessa gestão.

Podemos aqui incorporar algumas considerações na proposta do PL de regulamentação de inteligência artificial para que nós evitemos algumas causas danosas aos cidadãos. Uma delas é que nós possamos considerar os impactos dos sistemas com uma supervisão de determinação humana, pois os humanos são éticos o suficiente e legalmente responsáveis por todas as fases dos ciclos dos projetos quando está sendo construído algum sistema de inteligência artificial.

A inclusão de gênero não deve reproduzir as desigualdades que nós encontramos no mundo real, e hoje nós vemos isso fielmente em vários sistemas. Antes de nós começarmos um projeto, devemos dar um passo atrás e perguntar se a decisão de nós automatizarmos um processo é a escolha certa – até respondendo já à pergunta que uma das pessoas fez –, porque automatizar um projeto, automatizar um processo pode tornar mais difícil para nós inspecionar e observar e testar em uma possível auditoria futura.

(Soa a campainha.)

O SR. WALTER MARINHO – Os projetos devem levar em consideração os impactos que podem ter no bem-estar das partes interessadas e comunidades afetadas, ou seja, o cidadão deve ser considerado em sua diversidade antes de implementar um sistema e uma forma de prever e evitar efeitos prejudiciais.

Outra recomendação é um princípio de não dano discriminatório, que ajuda a garantir que um sistema não tenha impactos discriminatórios ou injustos na vida das pessoas. Isso afeta aqueles com características protegidas ou conteúdos que incluam orientações sexuais, religiosas, políticas, deficiência, entre outras, as quais não deveriam constar.

Um exemplo – sei que volto a passar um pouquinho o tempo – é que um fornecedor propôs uma solução de aprendizado de máquina para admissão de universitários, que previa a característica dos alunos com melhor desempenho. A universidade, ao considerar esse impacto, viu que o *software* dessa universidade corria riscos diversos, inclusive de vieses de discriminação, ou mesmo diminuindo a diversidade.

Outro fato ocorrido na pandemia foi uma universidade deixar de submeter os exames em nível humano e deixar que uma inteligência artificial o fizesse por ela. O resultado foi um viés machista,



SENADO FEDERAL
Secretaria-Geral da Mesa

deixando todas as mulheres de fora do curso de Medicina. E, ao fazer uma auditoria, viu-se que essas mesmas mulheres tinham nota suficiente para entrar nesses cursos.

Eu tenho ainda algumas considerações para fazer, que vou fazer em cima das perguntas feitas, mas é muito importante que a ética e a governança estejam incrustadas no PL para que possam ser discutidas e que venhamos a mitigar ou retirar totalmente os vieses.

O impacto de um viés na vida de uma pessoa que pode ser condenada, como o Sr. Daniel colocou ali, pois, no caso, aqui, é um respaldo que os auxilia... Mas temos casos em estado dos Estados Unidos, estado americano, onde o juiz é uma inteligência artificial. Então, se ele vem com dados sujos, ele amplifica vieses racistas. E nós temos um *case* em que, em determinadas regiões onde havia um grande número de crimes, ele automaticamente sugeria que pessoas daquela região poderiam ter uma maior chance de cometer crimes. Automaticamente, ele condenava as pessoas daquela região com crimes maiores do que outras com penas iguais ou com crimes iguais de outras regiões. Ou seja, ele não tinha uma certa visão cognitiva, por isso a importância de uma revisão humana.

Então, ficam as minhas considerações iniciais, e estou à disposição para nós debatermos mais sobre o tema.

Obrigado.

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO) – Obrigado, Dr. Walter.

Passo a palavra agora ao Dr. Giovanni Cerri, Presidente do Conselho de Inovação do Hospital das Clínicas da Faculdade de Medicina da Universidade de São Paulo (Fmusp).

O SR. GIOVANNI CERRI (Para expor.) – Queria cumprimentar o Senador Eduardo Gomes, Presidente desta audiência, também agradecer ao Senador Marcos Pontes e cumprimentar os demais componentes da mesa e todos os presentes.

Eu considero esse assunto extremamente importante para a área da saúde e eu queria fazer um breve histórico em relação à nossa atuação na questão da inteligência artificial.

Em 2017, o Hospital das Clínicas e a Faculdade de Medicina da USP criaram o Inova HC, que é um *hub* de inovação aberto para promover iniciativas de inovação na área da saúde, tentando compensar o atraso que nós tivemos na questão de estimular soluções para a saúde que estivessem adequadas para o nosso ecossistema e que também pudessem ajudar a reduzir o custo de saúde no Brasil.

Em 2020, nós criamos no *hub* do Inova HC um laboratório de inteligência artificial que visa pesquisar e desenvolver soluções de inteligência artificial para a saúde, utilizando o Hospital das Clínicas, que é o maior centro hospitalar do país, para justamente testar essas soluções. Aliás, nessa inauguração, estiveram presentes o Senador Marcos Pontes e o Prof. Marcelo Morales, que perceberam a importância de o Brasil buscar soluções adequadas à nossa realidade e também com um banco de dados, que é um banco que representa a nossa realidade social, racial. E o Hospital das Clínicas representa bem isso, porque atende toda a população, a população do Brasil inteiro.

Então, o laboratório de inteligência artificial partiu de parcerias, parcerias como inovação se fazem com parcerias públicas e privadas. Tivemos apoio de diversas empresas. A primeira parceria foi com a Siemens, e depois temos hoje um trabalho importante com a AWS, com a Huawei, com diversas empresas nesse sentido de desenvolver soluções de inteligência artificial.

Destaco que o nosso primeiro projeto foi durante a pandemia, quando nós levantamos uma plataforma de inteligência artificial com que, com mais de 70 hospitais coligados à nossa plataforma, nós respondíamos, em poucos minutos, se o paciente, baseado na tomografia de tórax, tinha covid e qual a extensão da doença, num momento em que o exame laboratório demorava uma semana ou dez



SENADO FEDERAL
Secretaria-Geral da Mesa

dias para ficar pronto. Então, essa plataforma foi importantíssima no começo da pandemia para poder apontar quais os casos graves que precisavam de uma atuação clínica.

E o nosso último projeto desenvolvido foi uma *startup* que identifica tumores de fígado em pacientes com doença crônica, com cirrose. Ou seja, o nosso laboratório já busca e já trabalha com essa missão de desenvolver inteligência artificial na saúde.

E qual seria esse potencial de inteligência artificial na saúde? Nesse potencial, inicialmente, nós temos que avaliar um pouco a questão dos investimentos, que é algo que envolve muitos bilhões em todo esse mercado global de saúde digital. Hoje nós temos mais de 1.380 *healthtechs* no Brasil que foram criadas nos últimos anos, e eu devo dizer que muitas dessas *healthtechs* desenvolvem ou estão vinculadas a projetos de inteligência artificial. Ou seja, para que nós ajudemos a impulsionar toda essa solução em saúde para o nosso mercado, para a nossa realidade, é muito importante que essas *healthtechs* sejam apoiadas tanto do ponto de vista financeiro como de facilidades para que elas possam se desenvolver.

O próprio Ministério da Saúde criou uma Secretaria de Informação e Saúde Digital – até foi por nossa recomendação –, na qual se colocam algumas prioridades. A primeira é a informatização nos diferentes níveis de atenção, a interoperabilidade de dados, que é uma coisa essencial para o desenvolvimento da saúde do país, e todo esse fomento ao sistema de inovação, no qual a inteligência artificial é uma das áreas prioritárias.

Quando nós observamos o potencial de aplicações na cadeia da saúde, nós caracterizamos que isso vai desde os novos produtos, que são ferramentas que auxiliam a pesquisa e o desenvolvimento de novos produtos, identificamos a inteligência artificial como fundamental para poder auxiliar a pesquisa de novos medicamentos e dispositivos médicos e é um grande apoio à pesquisa clínica. Hoje já existem algoritmos que possibilitam simular pacientes sem utilizar os dados dos pacientes reais, ou seja, pesquisas clínicas baseadas em dados existentes, mas que preservam a identidade e os dados dos pacientes reais.

Na questão da gestão populacional, para nós conseguirmos definir políticas de saúde, é fundamental que nós possamos analisar os dados. A inteligência artificial consegue ajudar os gestores a analisar esses dados com muito mais eficiência, dados em muito maior volume, e isso ajuda a aprimorar as políticas públicas.

Na questão da gestão e administração do complexo saúde, a inteligência artificial hoje já ajuda nas compras dos processos de licitação, ajuda na identificação de informações que possam dar ao gestor o caminho para decisões certas, ou seja, ajuda em todo o processo decisório na questão da gestão em saúde.

Em sistemas hospitalares, eu queria dar dois exemplos, pois, hoje, o médico e o profissional de saúde perdem muito tempo transcrevendo dados para o prontuário eletrônico. Na inteligência artificial generativa, já existe um algoritmo nesse sentido, que permite sumarizar a conversa do médico com o paciente ou as decisões clínicas, transcreve para o prontuário, e o médico avalia e assina, ou seja, pode ajudar, e muito, na redução do trabalho burocrático em hospitais, liberando os médicos e profissionais de saúde para as tarefas mais nobres.

Na assistência do paciente, são inúmeras as aplicações, diagnóstico e tratamento, na questão educacional e na aderência dos tratamentos.

Eu queria dar dois exemplos na área de radiologia, pois sou professor da área de radiologia. Nós utilizamos algoritmos de inteligência artificial, por exemplo, para identificar hemorragia intracraniana em pacientes com acidente vascular cerebral. Muitas dessas hemorragias intracranianas são mínimas,



SENADO FEDERAL
Secretaria-Geral da Mesa

são pacientes que entram na urgência e, às vezes, o radiologista não é o especialista, e a inteligência artificial ajuda para que esses casos não sejam perdidos e o diagnóstico não deixe de ser feito, porque, em hemorragia intracraniana, o tratamento, quanto mais rápido, mais eficiente e melhor benefício traz para o paciente.

E há também o amplo uso de inteligência artificial no diagnóstico de diversos tumores, também como auxílio na radiologia, como tumores de mama, e o nosso algoritmo, que nós desenvolvemos agora, recentemente, para tumores de fígado.

E aí vale um comentário que eu considero muito importante: hoje nós temos mais de 690 dispositivos médicos que utilizam a inteligência artificial, autorizados pelo FDA, até julho de 2023, em diversas especialidades. Hoje a maior aplicação ainda é radiologia.

Vale uma história: em 2019, houve uma previsão de um cientista de que o radiologista desapareceria porque a inteligência artificial iria substituir o radiologista.

(Soa a campanha.)

O SR. GIOVANNI CERRI – Hoje eu diria que o radiologista que não usa inteligência artificial será substituído pelo radiologista que usa a inteligência artificial. A inteligência artificial é uma ferramenta de trabalho fundamental para diversas áreas da saúde e ajuda a proteção e a segurança do paciente.

Devo destacar que isso já passa por uma grande fase de comprovação. Existem 50 estudos relevantes que avaliaram a implementação e a aplicação da inteligência artificial na prática clínica. Dentro desses estudos, o que se mostrou? Que 92% reportaram uma *performance* aceitável desses sistemas em comparação com a *performance* humana, ou seja, são complementares à atuação humana; 81% reportaram resultados positivos para a equipe médica; 78% reportaram resultados positivos para os pacientes.

Eu quero destacar que isso é o começo. A tendência da utilização da inteligência artificial vai ser aprimorar ainda mais esses estudos, mas é importante destacar que a inteligência artificial é como ferramenta e que a palavra final é do médico, do profissional de saúde, que tem que avaliar e tem que julgar o auxílio que a inteligência artificial dá para todo esse processo decisório.

E aí vamos um pouco para as questões e desafios regulatórios, que é um pouco o tema desta audiência.

Quando nós olhamos as aplicações – e nós temos aí presentes aplicações na área da saúde –, a pergunta é: todas as aplicações na área de saúde seriam de alto risco? E nós observamos que a inteligência artificial entra na pesquisa de desenvolvimento, no apoio à pesquisa clínica. Vários equipamentos de radiologia já incorporaram algoritmos de inteligência artificial – já incorporaram. Ajuda na alocação de vagas e gestão de filas. Quando se fala em gestão de filas, por exemplo, é no processo que ajuda na emergência. Na emergência, a seleção é por gravidade. Então, a inteligência artificial pode processar rapidamente as informações para que o paciente mais grave seja atendido prioritariamente. Hoje muitos pacientes acabam não sendo atendidos de forma adequada na urgência. A gestão de equipe de saúde, o auxílio ao diagnóstico. A inteligência artificial é muito importante na telemedicina, podendo orientar pacientes e ajudar no aumento da eficiência do setor saúde, no auxílio ao tratamento, na aderência ao tratamento. Ou seja, hoje a inteligência artificial é muito importante na melhoria da eficiência na saúde e cada vez mais será importante no aumento da segurança do paciente.

Por isso, eu diria que a cadeia de saúde é muito ampla. As aplicações da inteligência artificial no setor podem acontecer em etapas muito diversas, com diferentes níveis de risco e exposição ao paciente. Então, não se pode dizer certamente que todas as aplicações na área de saúde são de alto



SENADO FEDERAL
Secretaria-Geral da Mesa

risco, de forma nenhuma. Muitas aplicações são essenciais para a segurança do paciente e para o aumento da eficiência do setor, o que é necessário, em razão da questão de custo.

Daí eu derivo um pouco para a questão de que nós temos, na saúde, uma agência reguladora, que, aliás, funciona muito bem e que, na questão dos dispositivos médicos, já classifica o risco intrínseco desses dispositivos em baixo risco, médio risco, alto risco e máximo risco. Ou seja, a Anvisa já tem uma previsão de classificação de risco de novos dispositivos na saúde, enquanto o PL 2.338, no art. 17, considera todas as aplicações na área da saúde como alto risco. Eu vejo isso como um aspecto extremamente negativo.

Se nós considerarmos todas as aplicações da saúde de alto risco, nós vamos inviabilizar grande parte da pesquisa e da aplicação da inteligência artificial na saúde, vamos retardar o processo e vamos prejudicar a segurança do paciente em muitas situações em que ela já é utilizada.

Finalmente, como conclusão, eu diria que as aplicações do setor são muito diversas – gestão, educação, pesquisa clínica – para que sejam todas consideradas de alto risco.

A própria Anvisa discrimina níveis de risco variáveis para tecnologias da saúde, a depender de sua aplicação. O setor da saúde já é regulado e conta com longo de ciclo de desenvolvimento.

O PL, ao considerar todas as aplicações de alto risco, trará custos adicionais e obstáculos ao setor.

Ou seja, estas três reflexões eu faria: a importância da aplicação da inteligência artificial da saúde; não considerar todas as aplicações de alto risco; e considerar que a Anvisa é uma agência reguladora que já tem a preocupação com a questão dos riscos na saúde.

Por isso, como conclusão, eu diria que a redação do PL está genérica e não descreve como o setor funciona na prática.

Acreditamos que a categorização de todas as aplicações de inteligência artificial no setor de saúde como alto risco é um equívoco.

Então, esses são os comentários que eu queria fazer.

Agradeço a oportunidade de falar nesta audiência.

Muito obrigado, Senador.

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

Quero fazer duas observações.

Temos alguns Senadores conectados – o Senador Marcos do Val, o Senador Izalci – nesta sessão.

Observo que, sobre o que nós estamos falando, hoje, a *Folha de S.Paulo* traz matéria falando que, segundo o Dicionário Collins, "inteligência artificial" é a palavra do ano de 2023, tamanha a importância do tema.

Passamos, neste momento, a palavra à Crisleine Barboza Yamaji, Professora do Instituto Brasileiro de Mercado de Capitais (Ibmec), que está fazendo a sua participação remota nesta sessão.

Tem V. Sa. a palavra.

A SRA. CRISLEINE BARBOZA YAMAJI (Para expor. *Por videoconferência.*) – Muito obrigada.

Antes de tudo, gostaria de agradecer ao Senado e aos Senadores nas pessoas do Presidente Eduardo Gomes e do Senador Marcos Pontes, que fez, gentilmente, este convite.

Saúdo também todos os Senadores e membros aqui presentes.

Eu acho que o debate que está sendo construído é bastante relevante. Eu venho como professora de uma universidade bastante focada em inovação, tecnologia e setores produtivos; e sou advogada do setor financeiro e do setor produtivo.



SENADO FEDERAL
Secretaria-Geral da Mesa

Aqui, começando a trazer uma ponderação sobre o PL 2.338, que é bastante relevante, a gente, quando dialoga com os setores produtivo e financeiro, tem ouvido bastante que esse é um PL feito na prancheta, um PL de juristas para juristas e um PL produto da importação de um modelo regulatório que não se encaixa a nossa realidade nacional. Isso é o que a gente ouve; e ouve também a necessidade de descer para o chão de fábrica. Eu acho que este debate público impulsionado pelo Senado é bastante importante para complementar todas as discussões que vêm sendo feitas.

Nesse contexto, o que os setores produtivos sentem, antes de tudo, é um projeto que vem substituir projetos anteriores, que eram projetos bastante principiológicos, bastante flexíveis, bastante abertos à inovação e à tecnologia, e que vinham inseridos num contexto em que a nossa própria Constituição, quando traz um capítulo de ciência, tecnologia e inovação – ponderada, é claro, com a salvaguarda dos direitos fundamentais e essenciais na nossa Constituição –, prezava pelo fomento, pela promoção do desenvolvimento científico e tecnológico do país.

É um país que luta para se desenvolver, para ter presença de indústria nacional, para ter um setor financeiro forte e um cenário competitivo no âmbito internacional; um país que depende de um estímulo, de uma cooperação entre entes públicos e entes privados, um incentivo do mercado interno num cenário em que os recursos são escassos.

Diante do PL 2.338, a gente deixa de olhar os preceitos constitucionais para ciência, tecnologia e inovação.

Podemos dizer que ele não é exatamente um projeto de juristas para juristas, porque ele também, do lado dos juristas, deixa de olhar o cenário do direito econômico, que soma esforços do direito público, do direito privado e da legislação como um todo para impulsionar o desenvolvimento do país. Nesse sentido também, deixa de olhar a Estratégia Brasileira de Inteligência Artificial, tão bem colocada pelo Ministério de Ciência e Tecnologia, que trazia diversos dos preceitos constitucionais, que a gente está aqui colocando, claro, sem prejuízo da salvaguarda dos direitos fundamentais; trazia a questão do desenvolvimento sustentável, do crescimento inclusivo e do uso responsável da tecnologia. É disso que a gente está falando, mas, quando a gente pressupõe que praticamente todo o setor produtivo está dentro de alto risco, a gente não está prezando pelo uso responsável da tecnologia e a gente não está compreendendo efetivamente os benefícios e os impactos dessa inteligência artificial. E duas são as consequências.

Quando a gente traz aquele pacote de classificação de riscos importado do continente europeu, que não olha o cenário nacional, mas que coloca todo o setor produtivo em risco alto, a gente está pressupondo que o uso da tecnologia, o uso da inteligência artificial é negativo, ele é algo de que a gente quer se proteger a qualquer custo. Tanto é que, na linguagem, a gente vê que há uma grande preocupação com direitos de pessoas afetadas, mas todo direito tem um contraposto do dever. Onde está o dever, onde está essa ponderação um pouco mais equilibrada do setor produtivo e do setor financeiro, que já usam a inteligência artificial? Foi muito bem colocado aqui pelos colegas que a inteligência artificial não é algo novo, ela já vem sendo usada para a experiência do cliente, para a maior eficiência na produção, para uma categorização adequada dos dados, que sofrem, sim, de vieses, mas que são vieses humanos, que não são neutros. Então, há necessidade de uma supervisão, mas também cuidado, para se evitar pensar que o humano vai interferir em toda e qualquer mínima etapa do processo, porque, senão, não é uso de inteligência artificial.

Então, uma supervisão adequada e uma consideração dos setores produtivos que já têm uma regulação para isso... O setor financeiro é robusto em regulações para riscos operacionais e riscos sistêmicos, e o Banco Central tem pensado e refletido sobre o aprimoramento dessas normas diante do



SENADO FEDERAL
Secretaria-Geral da Mesa

desenvolvimento tecnológico, mas se colocando como um órgão regulador de incentivo e de fomento, e outros órgãos reguladores do país também. O Ministério de Ciência e Tecnologia tem tomado a frente para articular essas ações e olhar uma política nacional.

Então, é preciso evitar que o projeto de lei que saia seja um projeto de lei teórico, de implantação de transplante de modelos e, mais do que isso, que deixe de considerar que ele hoje traz diversas sobreposições normativas, porque se sobrepõe à Lei Geral de Proteção de Dados, à legislação consumerista, às questões de segurança cibernética, a questões, como foram colocadas aqui, de saúde, a questões trabalhistas, do desenvolvimento da pesquisa e da formação da nossa força de trabalho que vai lidar com essa inteligência, mas que já deve, diante desse PL, pressupor que é algo ruim, algo que faz mal para o país, em um contexto internacional em que um país deveria se posicionar para ganhar independência.

O Brasil é sedento por tecnologia, a gente tem uma inteligência nacional bastante forte, que usa tecnologia, e a gente tem posição de ponta, se bem fomentada e bem pesquisada.

Outras questões importantes nesse contexto é considerar que talvez – acho que isso já foi dito – esse projeto deveria substituir trazendo o espírito dos projetos anteriores, então, sendo principiológico, com os seus fundamentos, olhando diretrizes éticas que são inegociáveis e salvaguardas de direitos fundamentais, mas, ao mesmo tempo, fazendo uma total revisão do art. 5º em diante.

A parte de direitos tem que vir ponderada com a Estratégia Brasileira de Inteligência Artificial, olhando os setores que usam a inteligência e que produzem a inteligência pelo país. Toda a parte de classificação de riscos tem que ser absolutamente revista, olhando-se a questão talvez com enfoque naquilo que é inaceitável, que é excessivo, que afronta direitos fundamentais, mas permitindo a flexibilidade de uma abordagem baseada em riscos do setor financeiro e dos setores produtivos, para que eles próprios, estando inseridos no chão de fábrica, possam categorizar e estabelecer monitoramentos sem burocratizar demais.

Quando a gente olha as questões de direito, às vezes, a gente tem a impressão de que se quer documentar 100% de processos de uma forma bastante onerosa. Devem, sim, ser documentados, devem, sim, ser supervisionados, mas sob medida para cada setor, e isso não está sendo considerado.

Também nas governanças, a gente já tem padrões de proteção de dados pessoais pela própria ANPD (Autoridade Nacional de Proteção de Dados), que olha os dados pessoais, que são fundamentais. Vamos dizer que as inteligências artificiais, hoje, para 80%, 90% dos processos produtivos, usam dados pessoais e já têm uma autoridade alocada para isso, com uma regulação própria, que vai poder tratar, junto com a própria LGPD, com todo o tratamento do setor, dessas questões, fora os outros órgãos reguladores, fora o Ministério de Ciência e Tecnologia.

Então, eu venho aqui, enquanto professora e advogada, trazer essas preocupações, esperando que o PL 2.338, que foi um excelente trabalho, seja revisto nesses aspectos, que seja focado numa regulamentação principiológica, flexível e que considere a nossa posição e a nossa necessidade de desenvolvimento do país, um desenvolvimento inclusivo, um desenvolvimento tecnológico responsável, mas não algo que seja pressuposto como deletério, como ruim, como uma afetação negativa para a população brasileira, que vem inserida em um contexto internacional e quer ter sua independência, sua produtividade, sua pesquisa e também sua posição de ponta no cenário internacional.

Agradeço muito o espaço e o convite do Senado e fico à disposição para os debates e para contribuições adicionais.

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO) – Muito obrigado.



SENADO FEDERAL
Secretaria-Geral da Mesa

Passo, neste momento, a palavra ao Dr. Marco Antonio Lauria, membro do Conselho da Associação Internacional de Inteligência Artificial.

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

O SR. MARCO ANTONIO LAURIA (Para expor. *Por videoconferência.*) – Muito obrigado.

Um bom dia a todos!

É um prazer estar aqui participando desta 14ª sessão da Comissão de IA do Senado. Meus agradecimentos ao Presidente da Comissão, Senador Eduardo Gomes, ao Senador Astronauta Marcos Pontes, aos demais Senadores e demais presentes nesta sessão.

Eu trabalhei por muitos anos na IBM, por mais de 30 anos, eu fui responsável pela área de soluções cognitivas, que era a área responsável pelas soluções com Watson, a inteligência artificial para a América Latina. Nós chegamos a implementar mais de 600 projetos não só no Brasil. Com a América Latina, tínhamos um time de, mais ou menos, 1,5 mil pessoas trabalhando com isso. Depois que eu saí da IBM, fui um dos fundadores da Associação A2iA, uma associação que visa promover o uso racional, o uso responsável de inteligência artificial no Brasil, reconhecendo que isso é uma tecnologia fundamental para o desenvolvimento do país.

Acho o trabalho desta Comissão importantíssimo. É um momento em que este tema está sendo discutido no mundo inteiro. Só nesta semana, tivemos uma *executive order* do Presidente Biden nos Estados Unidos, o código de conduta pelo G7, o *International Code of Conduct for Advanced AI*, e os dois saíram na segunda-feira. Hoje mesmo, está acontecendo o *AI Safety Summit*, no Reino Unido, ou seja, o mundo inteiro está discutindo esse tema no momento, o que mostra como isso é importante.

Eu queria trazer alguns dados aqui mostrando uma agenda positiva, como já foi falado aqui, foi falado pelo Bruno, foi falado por outros. O primeiro documento, um documento já antigo, que falava do impacto da inteligência artificial e tentava olhar qual o impacto disso no PIB mundial foi um documento da Pricewaterhouse de 2017, chamava *Sizing the Price*, como se fosse Dimensionando o Prêmio, e estimava um crescimento do PIB global devido à IA, até 2030, em 14%. Ou seja, a IA traria um crescimento adicional de 14% do PIB, o que equivaleria, mais ou menos, a US\$16 trilhões, que é o que está-se falando. Mas esses 14% não aconteceriam de forma igual no mundo inteiro. O estudo previa que Estados Unidos e China seriam os dois que mais se beneficiariam disso: a China com 26% do PIB e os Estados Unidos com 14% do PIB; a Europa e a Ásia desenvolvida, na casa de 10% do PIB; e a América Latina, incluindo o Brasil, a África e a Oceania, na faixa de 5% a 6% do PIB. Ou seja, esse primeiro estudo, que é de seis anos atrás já, estimava um crescimento do Brasil que era metade do crescimento da Europa, um terço do crescimento americano e quase um quinto do crescimento que a China iria ter com a inteligência artificial.

Depois desse estudo, obviamente, teve outros estudos mais recentes. Tem vários estudos de 2023 aí feitos também, pela McKinsey, pelo Bank of America, Goldman Sachs. Eles estimam que o impacto pode ser até maior: com o surgimento aí da IA generativa, que foi mencionado aqui, nesse último ano, esse ganho de produtividade eles estão dizendo que pode ser mais: de 15% a 40% de ganho adicional de produtividade, o que levaria esses números a se tornarem números mais relevantes ainda. E, se nós olharmos em retrospecto o que aconteceu na primeira e na segunda revolução industrial, com o advento das máquinas a vapor, da máquina à combustão e da eletricidade, e a gente olha que a inteligência artificial talvez seja o principal fator dessa quarta revolução industrial. Deixar de capturar esse benefício, ou seja, não aproveitar essa onda certamente seria um erro, iria levar a um aumento da desigualdade, iria levar a uma perda da competitividade por parte do país. A IA atinge todos os setores, acho que a gente teve pouquíssimas tecnologias que tiveram a abrangência que a IA tem. Ela vai impactar o



SENADO FEDERAL
Secretaria-Geral da Mesa

trabalho, o futuro do trabalho, ela impacta o lazer, impacta aquilo a que a gente assiste, a informação que a gente consome, ou seja, é importantíssimo aproveitar essa onda e colher os benefícios da inteligência artificial.

E aí eu me remeto a um estudo que nós fizemos no comitê de inovação – é um dos comitês da I2AI –, tem mais de 50 membros que participaram desse estudo, e nós focamos quatro setores, três dos quais são objetos da discussão aqui: o agronegócio, o setor de indústria e o setor de finanças seguras; a gente olhou outras coisas, olhou logística e transporte, olhou algumas outras coisas. Mas, na época – e esse estudo tem mais de um ano –, foram identificadas mais de mil *startups* brasileiras – eu não estou falando de *startups* internacionais –, mais de mil *startups* brasileiras que usavam de alguma forma inteligência artificial. Se nós fizermos essa análise hoje, provavelmente o número vai ser... certamente o número vai ser maior. Eu acredito que a gente deva ter umas 1,4 mil, 1,5 mil *startups* pelo menos, e todas elas estão incluindo a inteligência artificial, ou seja, é uma questão de tempo até que esse número se torne um percentual maior ainda. E nós mapeamos as principais funcionalidades em alguns mapas mentais, documentamos na época mais de 300 casos de uso, e eu queria mencionar alguns deles aqui.

Por exemplo, na parte da agricultura – que é importantíssima para o país, representa mais de 20% do PIB –, a detecção antecipada de pragas e o controle de pragas permitem a aplicação seletiva de defensivos agrícolas. A análise do solo e a previsão climática permitem um uso mais racional da água. Essa gestão e mapeamento do solo permitem um aplicação racional de fertilizantes e corretivos, ou seja, nós estamos falando de uma agenda de aumento de produtividade de alimentos de um setor que é essencial para o país, mas também uma agenda totalmente alinhada com ESG e com a preocupação climática que a gente tem hoje, com menor uso de defensivos, menor impacto ao planeta, uso mais racional de águas. E tem uma série de outros fatores: a recomendação para manejo e colheita, para otimização da produção; toda a parte de otimização de rotas, com a parte de gerência de frotas, gerando economia de combustível, na parte de logística e transporte.

O Bruno falou de indústria, uma indústria que depende desses benefícios e dessa tecnologia, então a gente tem aí redução e prevenção de acidente por rastreamento de pessoas e máquinas, detecção antecipada de falhas, melhora de qualidade, treinamentos usando realidade virtual, realidade aumentada, otimização da produção com automação; na parte de engenharia, simulações, engenheiros digitais; em finanças, como foi falado aqui, detecção, prevenção de fraudes, ou seja, tem uma série de exemplos conhecidos e documentados no uso de inteligência artificial e por isso é importantíssimo a gente capturar esses exemplos.

Então eu diria que a chave aqui é um equilíbrio entre o incentivo e o controle, um balanço entre oportunidade e risco. Como é que a gente captura a oportunidade mitigando os riscos? Se nós só focarmos nos riscos, nós não vamos aproveitar a oportunidade. Eu entendo que o principal objetivo do PL não é o foco da oportunidade, aqui foi mencionado também da Ebia, tem a PEC 31, que fala de um plano de investimento, mas, se a gente quiser capturar essa oportunidade, surfar essa onda, vamos precisar criar um ecossistema de inteligência artificial, são várias entidades aqui, ter um plano de investimento, investir em pesquisa e desenvolvimento. Qual é a vocação que o Brasil quer ter em inteligência artificial? O que ele vai fazer? Nós vamos desenvolver modelos, vamos ser só um desenvolvedor? O que o Brasil quer fazer? Nós precisamos fomentar a inovação, precisamos ter programas para a educação e a capacitação. E, quando a gente olha o PL, embora ele fale de medidas para fomentar a inovação, a única coisa que aparece, na Seção III do Capítulo VIII, é uma leve menção falando de um *sandbox* regulatório, ou seja, definitivamente se fala muito mais dos riscos do que das oportunidades, e tem essa menção pequena.



SENADO FEDERAL
Secretaria-Geral da Mesa

Eu queria fechar, no pouco tempo que falta, falando de alguns pontos de atenção, alguns deles que já foram colocados aqui. Quando a gente fala de categorização de risco, a maioria dos países adota uma matriz em que entra severidade e probabilidade de dano, isso vem desde o modelo de Singapura. Aqui tem uma lista imensa de atividades que foram classificadas como de alto risco, setores inteiros em que existem aplicações, como foi demonstrado aqui, no setor de saúde ou de educação, que não representam risco, ou seja, o risco não pode ser uma coisa setorial, mas tem que estar focado aqui nas aplicações.

Quando a gente fala de responsabilidade civil, aqui menciona apenas duas entidades, o fornecedor e o operador, mas a realidade é muito mais complexa do que isso. Esses modelos são desenvolvidos em camadas, tem os grandes modelos de linguagem, tem modelos abertos, tem uma série de coisas em consideração e no fundo nós vamos ter muito mais entidades do que as duas mencionadas nesses itens, em alguns casos, o risco é muito maior do que está relacionado à inteligência artificial. Quando nós falamos, por exemplo, de carros autônomos, o que aparece aqui, ou quando nós falamos de uso médico, alguns usos dependem de sensores, de atuadores, de robótica, ou seja, o risco tem que ser mais abrangente e incluir esses casos.

Finalmente, quando a gente fala de supervisão e fiscalização, a recomendação é que se deveria aproveitar as entidades que já existem hoje. Hoje se falou da Anvisa, tem entidades que regem o setor de finanças e os demais, ou seja, deveria existir uma coordenação, mas o aproveitamento dessas entidades que, de alguma forma, já gerenciam esses setores individuais e provavelmente são as mais indicadas para olhar o uso de inteligência artificial para esses determinados setores.

Então, com isso eu me coloco à disposição para continuar esse debate, agradeço a chance de estar aqui. E, se puder deixar uma mensagem, é a busca desse equilíbrio. Acho que existe uma oportunidade enorme, o Brasil não pode ficar de fora. E a gente tem que buscar um equilíbrio entre a oportunidade e o risco.

Muito obrigado aí pela participação, pela chance.

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO) – Muito obrigado.

Nesse momento, a gente passa a palavra ao Dr. Robert Janssen, Presidente da Associação das Empresas Brasileiras de Tecnologia da Informação do Estado do Rio de Janeiro.

O SR. ROBERT JANSSEN (Para expor. *Por videoconferência.*) – Uma boa tarde para todos. Eu saúdo o Presidente, Senador Eduardo Gomes. Saudações também para o Senador Astronauta Marcos Pontes e todos os demais Senadores presentes, aos colegas aqui também presentes e a todo mundo que nos ouve.

É uma oportunidade singular estar neste momento da história mundial, não é só brasileira, endereçando o que pode se tornar o grande diferencial de você poder proporcionar prosperidade para as nossas nações.

Sem dúvida nenhuma, é um desafio enorme buscar o equilíbrio que o colega Marco Lauria traz. A gente entende que existe o dever da proteção e a salvaguarda dos direitos fundamentais da sociedade, mas existe o dever também de proporcionar desenvolvimento socioeconômico e competitividade do Brasil no contexto mundial.

Nós, enquanto entidade empresarial do setor de tecnologia, representamos também dentro de uma entidade mundial o Brasil, chamada WITSA, The World Innovation, Technology and Services Alliance, que recentemente fez uma publicação a partir de um estudo de como é que ele gostaria de poder recomendar porque diversos governos ao redor do planeta deveriam estar buscando o equilíbrio



SENADO FEDERAL
Secretaria-Geral da Mesa

entre a salvaguarda dos direitos fundamentais e o fomento da prosperidade através de um grande tsunami tecnológico chamado inteligência artificial.

Eu gostaria que pudessem projetar aqui a tela da nossa apresentação. A gente gostaria de poder compartilhar o que foram os pontos principais destacados nesse estudo porque depois a gente inclusive gostaria de fazer com que ele estivesse disponível para todos.

A minha apresentação, por favor, caso ela esteja disponível. Então, um instante só.

Rapidamente, a Assespro é uma entidade que tem o propósito de promover o desenvolvimento da sociedade com aplicação de inovação e tecnologia. A gente trabalha para que a gente seja sempre um vetor estratégico nesse processo e que a gente dê prosperidade para toda a sociedade.

Estamos desde 1976 articulando essa divisão, esse propósito, com mais de 2.500 empresas em torno de 20 estados. A WITSA é a organização mundial que eu mencionei. E esse documento que a gente produziu busca justamente fazer uma recomendação que ajuda a buscar o equilíbrio que é necessário entre a salvaguarda da proteção dos direitos fundamentais junto com a promoção e o fomento de prosperidade através do que nós estamos chamando desse tsunami tecnológico, que é a inteligência artificial.

E uma das principais conclusões é que, como já foi colocado pelos meus colegas, diversos setores já têm dentro de si uma normatização e um *oversight*, um controle sobre como atuam com relação ao risco. Então, os próprios setores já estão espontaneamente criando as suas normatizações voluntárias e o mercado todo adotando isso de uma forma sistêmica.

E o que apareceu nesse estudo? Foi que nós devemos estar buscando, sim, os *frameworks* éticos e, ao mesmo tempo, um desenvolvimento sustentável. Tem que se verificar sempre a inclusão e, ao mesmo tempo, a privacidade e a segurança de dados.

Aqui eu queria fazer um parêntese específico sobre inclusão, que traz um tema chamado diversidade. Em 2006, foi feito um estudo sobre a miscigenação das raças em nível mundial. E a conclusão a que se chegou é que a raça brasileira é a mais miscigenada do planeta. Consequentemente é a que melhor contém o contexto da diversidade. Eu chamo a atenção para que o Brasil pode ser o protagonista nesse processo de ter a inserção do componente de diversidade no contexto da inteligência artificial.

A mitigação de preconceito e a transparência também são questões especificamente importantes, junto com a responsabilidade e a questão da colaboração.

Eu parableno novamente o movimento que o Senado faz. E obviamente sabemos que, no Congresso, existe o mesmo processo, mas novamente voltamos aqui ao desafio desse equilíbrio, que é fundamental.

Nós temos também a supervisão humana e a aprendizagem e melhoria contínua, e isso é um impacto global. O impacto global é quando todo mundo entende que esse tsunami tecnológico é o que chamamos de uma mudança de paradigma, uma mudança de patamar. No inglês, a gente chama de *game changer*.

E no final, o que a gente gosta de recomendar, gostaria de recomendar como política pública é que esses *frameworks* políticos e regulatórios devem usar o uso responsável, mas ao mesmo tempo fazer o desenvolvimento da inovação, como um fundo indutor para que haja mais desenvolvimento socioeconômico em todos os países.

Esse documento, *Construindo Confiança e Cumprindo a Promessa da Inteligência Artificial* está disponível. Eu vou deixar aqui um QR code para que todos possam baixá-lo.



SENADO FEDERAL
Secretaria-Geral da Mesa

E aí, os princípios que foram desenvolvidos a partir do estudo, que a gente sugere: que os governos devem evitar os preconceitos pró-humanos; devem regular o desempenho, e não os processos; regular os setores, e não tecnologias. A gente não quer colocar uma camisa de força no que vai ser a principal tecnologia que vai promover o desenvolvimento socioeconômico e gerar mais inclusão. Nós também queremos evitar a miopia da IA, definir a IA de uma forma precisa e restrita, para que assim possamos também fazer uma definição correta dos riscos. E o que já falaram aqui os meus colegas, já existem regras que podem ser usadas para poder fazer essa proteção e encaminhamento do que é a definição de um dano maior. E, ao mesmo tempo, não fazer discriminação, tratar as empresas igualmente; buscar a *expertise*; adotar, sim, *sandboxes* regulatórios; e adotar o uso governamental de IA de forma responsável. E, sim, deve-se fazer essa definição do risco, mas uma definição de risco que faça justamente o balanço entre o que é a proteção e a salvaguarda de direitos fundamentais e o que é que promove a prosperidade da sociedade como um todo.

Eu quero agradecer a oportunidade de estar aqui compartilhando esse estudo que foi feito pela WITSA, junto com a Assespro, em nível mundial. E quero deixar também aqui então o acesso, para que as pessoas possam baixar. E coloco-me à disposição na sequência. Muito obrigado.

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO) – Muito obrigado.

Agora vamos para a última exposição, do Daniel Stivelberg, que é coordenador do Grupo de Governança e Regulação de Dados da Zetta.

Eu vou conceder a palavra pelo prazo de dez minutos.

O SR. DANIEL STIVELBERG (Para expor.) – Obrigado, Presidente, Senador Eduardo Gomes, pela paciência de ouvir, assim, essa pletora de especialistas. A gente sabe que esta aqui não é a única audiência pública que aconteceu, que houve diversas audiências públicas que já aconteceram. A gente sabe que o processo é maçante, o de conformação da convicção para a elaboração e a produção normativa, mas é um grande privilégio estar aqui podendo auxiliar, dando subsídios para a criação de uma norma que a gente entende que é transformacional e que é infraestrutural para todos os setores. Queria agradecer também ao Senador Astronauta Marcos Pontes, obviamente, pelo convite.

A Zetta é formada por empresas de tecnologia financeira, basicamente. Hoje a gente tem crescido muito, são instituições reguladas e não reguladas, mas que atuam em diversas verticais, conformando, algumas delas, verdadeiras plataformas de transformação digital do setor financeiro. Então, tem empresas que podem atuar com criptoativos, *marketplaces* e contas transacionais, tem empresas que vão atuar especificamente numa vertical, mas a gente tem muitas delas que já estão crescendo, que já estão na América Latina. Algumas são empresas argentinas, mas, principalmente, a gente tem empresas brasileiras. Aqui, o que a gente tem é um ecossistema de produção de tecnologia e de contratação de trabalhadores superespecializados. O uso de tecnologias de transformação digital faz parte do modelo de negócios dessas empresas.

Então, quando a gente fala em regular a inteligência artificial, a gente tem que pensar especificamente no impacto que essas novas normas vão ter para esse tipo de ecossistema, um ecossistema que gera transformação econômica, mas que também tem impactos secundários muito positivos na economia como um todo, com empregos altamente especializados, profissionais com uma remuneração acima da média nacional e, obviamente, que vêm fazendo a transformação de um setor que outrora era um setor não tão dinâmico, fazendo aquele processo de inclusão financeira com redução de custos.

A gente tem, nesse ecossistema, mais de 100 milhões de clientes atendidos, uma economia com tarifas de mais de R\$60 bilhões nos últimos dez anos, o que, de fato, para nós é uma grande satisfação,



SENADO FEDERAL
Secretaria-Geral da Mesa

a de poder ser transparente e usar a tecnologia a serviço das pessoas para que elas possam ter acesso às suas finanças com mais autonomia, com mais conhecimento. A Zetta é movida por três pilares: competitividade, inovação e inclusão financeira.

Quando a gente fala de inteligência artificial, a gente está falando justamente do tronco de uma árvore. Essa metáfora, essa analogia, acho bastante relevante, Senador Eduardo Gomes, porque demonstra como essa tecnologia é transformacional, fundamental para a continuidade do crescimento desse setor e para a continuidade do bem-estar financeiro das pessoas.

A inteligência artificial e as suas diversas técnicas, juntamente com a abertura responsável de dados – e o Banco Central tem sido um grande parceiro nisso, na criação de arranjos regulados em que as pessoas possam exercer seus direitos fundamentais de privacidade e proteção de dados – tem conseguido mover um setor que outrora era mais opaco, menos transparente, para algo extremamente fácil e intuitivo de ser utilizado e cujas decisões sejam em favor dos interesses dos clientes, das pessoas, e não necessariamente em favor dos interesses de terceiros. Então, a gente vê a inteligência artificial como um tronco dessa grande árvore para a construção do sistema financeiro digital.

Esse estudo é um pouco antigo, mas mostra que as diversas técnicas de inteligência artificial vêm sendo aplicadas no setor, algumas delas já com necessidade de observar normas regulatórias setoriais.

Então, os principais casos de uso... E acho que a grande mensagem desse eslaide é: a gente usa a inteligência artificial para apoiar decisões de crédito ou de gestão de riscos de crédito. A gente usa a inteligência artificial para apoiar medidas de prevenção e combate à fraude. Existem aplicações de inteligência artificial utilizadas para melhorar a jornada dos clientes no seu autoatendimento e assim por diante e diversos outros processos operacionais e de gestão integrada de riscos.

O setor financeiro é um setor que se diferencia por causa da questão do risco sistêmico. Quando a gente faz intermediação financeira de dinheiro de uma parte para a outra, a gente trabalha com o que se chama de risco sistêmico e, em decorrência disso, existem reguladores especializados que vão buscar normatizar essas instituições, aquelas que são reguladas, para que elas tenham uma governança de modelos que atendam a essas normas setoriais.

E como essa regulação setorial induz essa governança? Esse é um modelo genérico de funcionamento das instituições, mas nem todas podem funcionar exatamente dessa forma – aqui é para fins de ilustração –, mas a gente tem um modelo de defesa de três linhas. A instituição, basicamente, a finalidade dela é criar produtos inovadores para facilitar a vida financeira das pessoas, por exemplo. E lá na primeira linha tem os times que vão criar esses produtos. Eles são donos de risco. Eles trabalham com os times de segunda linha, que são basicamente os times que estão observando em que medida o risco pode ou não ser assumido nas primeiras linhas e em que medida aquele produto está atendendo às obrigações regulatórias.

Nessa segunda linha, a gente tem grupos de conformidade, que em inglês eles chamam de *compliance*, a gente tem o DPO, que é o Data Protection Officer, o encarregado de dados. A gente tem um processo inteiro de gestão de riscos e rituais de gestão de riscos de modelos.

Quando a gente pensa numa instituição financeira e numa assemelhada, a gente fala de famílias de modelos. A gente tem famílias de modelos de crédito, de fraude, de atendimento ao cliente, até mesmo para arbitrar, por exemplo, capitais para evitar quebra de solvência. Então, todos esses modelos, cada qual com o seu tipo de risco. Não existe apenas um risco no sistema financeiro.

E aí, pensando um pouquinho na perspectiva das autoridades reguladoras, a gente tem, obviamente, a terceira linha, que é a auditoria interna, em que se pode auditar todo esse processo de criação de modelos e assim por diante.



SENADO FEDERAL
Secretaria-Geral da Mesa

Quando a gente pensa em regular a inteligência artificial, a gente pensa em três preocupações éticas fundamentais. A questão da qualidade do modelo, ou seja, a sua capacidade preditiva, a sua acurácia, se ele, de fato, vai dar *inputs* que representem aquela realidade. Você tem a preocupação relacionada ao impacto do modelo para as pessoas, que são preocupações normativas, de proteção de dados, de não discriminação ilícita e abusiva e assim por diante. E a gente tem uma preocupação com a responsabilidade e a rastreabilidade desses modelos.

São três grandes preocupações e essas três grandes preocupações já estão plasmadas no ordenamento jurídico brasileiro quando a gente fala para o sistema financeiro. A gente tem órgãos normatizadores e supervisores do Sistema Financeiro Nacional observando a qualidade dos nossos modelos, com a capacidade de auditá-los inclusive, com rastreabilidade de *datasets*, com auditoria sobre dados de treino e assim por diante.

Quando a gente fala de impactos para as pessoas, ou seja, de direitos fundamentais, se não existe uma discriminação ilícita ou abusiva, se uma pessoa pode pedir a revisão de uma decisão automatizada, dentre outras questões, a gente tem o Sistema Nacional de Proteção de Dados. E aqui também engloba, por exemplo, a Lei do Cadastro Positivo. Quando a gente vai olhar os dados que precisam entrar no modelo de crédito, por exemplo, Senador, a gente precisa olhar se esses dados estão atendendo à Lei do Cadastro Positivo.

Então, a gente tem um controle *ex ante* social sobre que dados entram no modelo de crédito, também isso estabelecido no ordenamento jurídico.

E a gente tem o Sistema Nacional de Defesa do Consumidor, que já estabelece regras de transparência, regras de qualidade dos produtos, o papel de fornecedor e consumidor e assim por diante. Então, a gente tem as competências fundantes de uso responsável de IA já estabelecidas para o setor financeiro e para entidades assemelhadas. E a gente gostaria que o Senado, se possível, levasse isso em conta na construção, na produção normativa, de modo que a gente não tenha colisão entre as normas e a gente possa ter deferência às regulações especializadas desse setor.

E, obviamente, a questão de ESG – diversidade, inclusão e ética –, são elementos que transcendem a própria norma. Não faz muito sentido a gente pegar elementos éticos e colocar numa lei porque deixa de ser ética para se tornar o direito; a ética é aquilo que supera, em que a gente vai além do direito. Então, tudo isso deve ser transversal numa governança de organizações como essa. Basicamente, essa é a pletora de organizações que já regulam e têm capacidade de auditar os nossos sistemas.

E, basicamente, para concluir, Senador, tentando ficar aqui nos dez minutos, a gente faz...

(Intervenção fora do microfone.)

O SR. DANIEL STIVELBERG – Obrigado, Senador.

A gente tem o seguinte pleito que eventual regulação não se dê sobre a tecnologia, mas sobre os usos dessa tecnologia.

(Soa a campainha.)

O SR. DANIEL STIVELBERG – Se um modelo de crédito é utilizado, se uma inteligência artificial é utilizada para prever a capacidade de adimplemento de obrigações financeiras, cabe ao Banco Central e aos órgãos normatizadores desse sistema financeiro olhar e entender se a qualidade do modelo atende aos requisitos epistêmicos daquele sistema.



SENADO FEDERAL
Secretaria-Geral da Mesa

Porque esse modelo não tem que observar tão somente direitos e garantias fundamentais – claro que tem –, mas existe um equilíbrio com o risco sistêmico, as normas de caráter prudencial do sistema financeiro. Então, esse ordenamento precisa cooperar. Essas autoridades, os órgãos do Sistema Financeiro Nacional e, no caso aqui, a Autoridade Nacional de Proteção de Dados, precisam cooperar para construir normas de uso responsável de modelos, eventualmente, para casos específicos como esse.

O marco legal precisa ser uma norma-parâmetro. Ele tem que ser um *framework* ou uma norma principiológica, como as pessoas vêm dizendo, porque as normas setoriais que vão trazer o conteúdo, sob o risco de a gente ter uma obsolescência normativa muito rápida, além, obviamente, de a gente não conseguir contemplar na lei todas as questões detalhísticas, prescritivas, que só uma autoridade setorial talvez teria condições e excluir mera automação.

Então, ali eu coloquei um conjunto de normas em vigor, especificamente sobre a questão de segurança e qualidade de modelos no sistema financeiro e em instituições financeiras e assemelhadas. Então, a gente tem elementos de ordem prudencial que a gente tem que observar nos nossos modelos. A gente tem a gestão integrada de riscos, que vai olhar o risco operacional, o risco de crédito, o risco de mercado, o risco social, o risco ecológico, porque a Resolução nº 4.557 foi recentemente atualizada pelo Banco Central. A gente tem normas de cibersegurança e contratação de serviços em nuvem.

No passado, não sei se o senhor se lembra, Senador Eduardo Gomes, mas falou-se na criação de uma lei de *cloud computing* no Congresso. A gente pensou: será que a gente precisa de uma lei de nuvem? Esse problema foi resolvido setorialmente. O Banco Central veio, sentou com o setor e estabeleceu normas de contratação de serviços relevantes em nuvem, porque o sistema financeiro precisa de um programa de continuidade de negócios; você não pode interromper eventualmente. O que o Banco Central fez? Sentou com os *players*, criou mecanismos de habilitação de *cloud providers* e lá está regulamentado. Então, a gente imagina que, no futuro, eventuais melhorias regulatórias possam se dar dessa forma.

E as normas em vigor, do ponto de vista de impacto para as pessoas, que é justamente privacidade, a LGPD e ANPD. Nossos modelos são intensivos em dados pessoais, mas já existem iniciativas de dados sintéticos em que processos computacionais simulam dados que não são dados pessoais e, através desses padrões, esses modelos aprendem. Então, é uma forma de melhorar a privacidade, processos de melhoria de privacidade. Regras de consumo, Código de Defesa do Consumidor, Lei do Cadastro Positivo. A gente tem obrigações civis, Código Civil; a questão da responsabilidade objetiva no Código de Defesa do Consumidor ou subjetiva imprópria, da LGPD, que é um tipo de responsabilidade muito especializado. Regras de KYC, isto é, "conheça o seu cliente". O Banco Central exige que as instituições, ao abrir uma conta de pagamento, por exemplo, conheçam, perfilizem o indivíduo, porque você precisa saber se a pessoa vai utilizar a conta para fins lícitos ou ilícitos, se ela vai lavar dinheiro e assim por diante. Então, a gente tem uma obrigação de perfilamento estabelecida na legislação.

E, finalizando, Senador, a questão dos direitos que a gente vem dizendo. A pessoa fala em princípios de uso responsável de inteligência artificial. Propósito benéfico, transparência e livre acesso, não discriminação, compreensividade, segurança, solidez e resiliência, que é a questão da robustez dos modelos, responsabilidade e prestação de contas, revisão e autonomia, todos esses princípios que estão na literatura de uma forma ou de outra, escritos de uma forma ou de outra, estão espalhados pelo ordenamento jurídico brasileiro de alguma forma. O que a gente precisa, eventualmente, é a criação de



SENADO FEDERAL
Secretaria-Geral da Mesa

um órgão, pensando aqui, que tenha a incumbência de fazer a conciliação e a harmonização desses recortes regulatórios e, eventualmente, dirimir conflitos.

Então, as duas recomendações, Senador, e aqui concluindo, são para, primeiro, fazer um juízo de necessidade de uma norma prescritiva, porque a gente entende que talvez seja o caso de uma norma, um marco legal como norma parâmetro; acolher as regulações setoriais e fazê-las prevalecer quando necessário para evitar a insegurança jurídica. E, finalmente, pensar na criação, talvez, de um comitê para dirimir eventuais conflitos de competência.

Senador, essa é a nossa contribuição. A gente fica à disposição, agradecendo novamente a atenção e a paciência das senhoras e senhores.

Obrigado.

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO) – Eu quero agradecer a todos os expositores.

Quero fazer duas considerações importantes para que não fique prejudicado esse debate, já que todas as exposições, inclusive das outras 72 apresentações, estarão e estão disponíveis na página da Comissão.

Entendo que... Vez ou outra, eu gosto sempre de resgatar o papel da Comissão de Juristas, que deu o primeiro passo e, por isso, tanto tempo depois a gente continua discutindo e ainda teremos a sessão de debates e a tramitação no Plenário do Senado e na Câmara dos Deputados. Era preciso dar esse primeiro passo. Por isso, eu falo sempre aqui do Ministro Cueva, da Dra. Laura, que foi a Relatora, de todos os membros da Comissão, e entendo que na conclusão dessa primeira etapa de debates, com essas 80 apresentações, nós iremos, nesta semana, intensificar o reconhecimento e a participação de todos os expositores.

Eu vou pedir, na próxima reunião da Comissão, ainda uma sessão adicional – eu estava conversando aqui com o Dr. Bandeira. Nós vamos fazer aqui uma audiência pública específica para todas as agências reguladoras, porque eu acho que está muito claro que a gente precisa levar essa contribuição para o debate do Plenário. E, em seguida, ainda propor, antes da sessão de debates... podemos inovar e conseguir a promoção de dois, três debates virtuais, com a participação de todos os expositores, para que a gente comece a chegar a um texto.

Eu sempre tive a aprovação de algumas PECs e de alguns projetos de lei... Então, a minha especialidade é tentar votar. Toda vez que algum Parlamentar fica muito envolvido tecnicamente com um tema é sinal de que ele vai para algum lado, porque tem outro muito forte disputando... O meu estado, o Estado do Tocantins, foi criado nas Disposições Transitórias da Constituição. Então, a característica do nosso estado é votar, buscar convergência e votar. E é isso que a gente vai propor com essa dinâmica. Vamos tentar inovar com esse debate.

Eu vou encerrar essa fase de audiências repetindo o que tenho dito aqui sempre: inteligência artificial é o primeiro assunto que eu reconheço, nesses 23 anos de mandato, em que você conversa com um especialista hoje, e, dois meses depois, quando você fala com ele, ele sabe menos.

E o Astronauta, o nosso querido Senador Marcos Pontes, ontem mesmo, trouxe uma série de dados sobre a influência do possível uso maléfico da inteligência artificial, contando a história lá do corte de energia em que, com apenas um movimento, tirou-se a energia de 250 mil pessoas. Aí há um momento que a gente pensa que o vírus é a gente – inteligência artificial.

Eu quero, com muita tranquilidade, com muita serenidade, dizer que o Presidente Rodrigo Pacheco priorizou o assunto, e é um assunto que – está claro aqui – a sociedade brasileira e mundial está discutindo. Então, a gente quer chegar a um bom termo ainda neste ano, com a análise dentro dos



SENADO FEDERAL
Secretaria-Geral da Mesa

120 dias, entendendo que o processo legislativo é assim: ele vai se aprimorando. O Congresso Nacional avança muito quando aprova boas leis e, na minha opinião, sempre avança muito mais quando deixa de aprovar leis ruins.

Com esse espírito, eu queria agradecer a todos e pedir à Secretaria da Comissão que a gente providencie para disponibilizar em QR code a apresentação de todos os expositores aqui para todos os expositores. E nós vamos tentar, nesses próximos dias, promover esse debate virtual para aprimorar e tentar chegar com o menor volume de problemas à fase pós-debate no Plenário do Senado Federal.

Eu quero agradecer a todos.

Nada mais havendo a tratar, declaro encerrada a presente sessão.

Muito obrigado. (*Palmas.*)

(Iniciada às 11 horas e 09 minutos, a reunião é encerrada às 12 horas e 56 minutos.)