



SENADO FEDERAL
Secretaria-Geral da Mesa

ATA DA 11ª REUNIÃO DA COMISSÃO TEMPORÁRIA INTERNA SOBRE INTELIGÊNCIA ARTIFICIAL NO BRASIL DA 1ª SESSÃO LEGISLATIVA ORDINÁRIA DA 57ª LEGISLATURA, REALIZADA EM 26 DE OUTUBRO DE 2023, QUINTA-FEIRA, NO SENADO FEDERAL, ANEXO II, ALA SENADOR ALEXANDRE COSTA, PLENÁRIO Nº 15.

Às dez horas e seis minutos do dia vinte e seis de outubro de dois mil e vinte e três, no Anexo II, Ala Senador Alexandre Costa, Plenário nº 15, sob a Presidência do Senador Astronauta Marcos Pontes, reúne-se a Comissão Temporária Interna sobre Inteligência Artificial no Brasil com a presença dos Senadores Izalci Lucas e Alan Rick. Deixam de comparecer os Senadores Carlos Viana, Styvenson Valentim, Veneziano Vital do Rêgo, Efraim Filho, Weverton (REQ 595 e 615/2023-CDir), Daniella Ribeiro, Vanderlan Cardoso, Nelsinho Trad (REQ 586, 587 e 588/2023-CDir), Fabiano Contarato, Chico Rodrigues e Laércio Oliveira. Havendo número regimental, declara-se aberta a reunião. A presidência submete à Comissão a dispensa da leitura e aprovação da ata da reunião anterior, que é aprovada e será publicada no Diário do Senado Federal. Passa-se à apreciação da pauta: **Audiência Pública Interativa**, atendendo ao requerimento REQ 4/2023 - CTIA, de autoria Senador Eduardo Gomes (PL/TO), com a finalidade de debater "aspectos centrais da regulação da Inteligência Artificial", a fim de tratar de abordagem principiológica, regulação baseada em riscos, regime de responsabilidade, governança multissetorial, estatuto de direitos, decisões automatizadas, supervisão humana, ética, privacidade e proteção de dados, com os participantes Rodrigo Badaró, Conselheiro Nacional do Ministério Público (CNMP); Ana Carla Bliacheriene, Professora da Universidade de São Paulo (USP); Leonardo Netto Parentoni, Professor da Universidade Federal de Minas Gerais (UFMG); Carlos Affonso de Souza, Consultor da Associação Brasileira de Internet (Abranet); Marcela Mattiuzo, Conselheira do Instituto Brasileiro de Estudos de Concorrência, Consumo e Comércio Internacional (Ibrac); Adriana Rollo, Líder da Comissão Especial de Regulação de Inteligência Artificial da Associação Internacional de Inteligência Artificial (A2IA); Cynthia Picolo, Diretora-Presidente do Laboratório de Políticas Públicas e Internet (Lapin); Fernanda Rodrigues, Coordenadora de Pesquisa do Instituto de Referência em Internet e Sociedade (Iris); Patrícia Peck, Presidente do Instituto Istart de Ética e Cidadania Digital; e Estela Aranha, Assessora Especial de Direitos Digitais do Ministério da Justiça e Segurança Pública (MJSP). O Senador Izalci Lucas faz uso da palavra. Nada mais havendo a tratar, encerra-se a reunião às doze horas e cinquenta e um minutos. Após aprovação, a presente Ata será assinada pelo Senhor Presidente e publicada no Diário do Senado Federal.

Senador Astronauta Marcos Pontes

Vice-Presidente 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/multimedia/evento/117871>



SENADO FEDERAL
Secretaria-Geral da Mesa

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO. Fala da Presidência.)
– Declaro aberta a 11ª Reunião da Comissão Temporária Interna sobre Inteligência Artificial no Brasil, criada pelo Requerimento nº 722, de 2023, com a finalidade de, no prazo de até 120 dias, examinar os projetos concernentes ao relatório final aprovado pela Comissão de Juristas responsável por subsidiar a elaboração do substitutivo sobre inteligência artificial no Brasil, criada pelo Ato do Presidente do Senado Federal nº 4, de 2022, bem como eventuais novos projetos que disciplinem a matéria.

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

As Sras. e os Srs. Senadores que concordam permaneçam como se encontram. (*Pausa.*)

Aprovada.

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

Pauta.

Informo que esta reunião se destina à realização de audiência pública com o objetivo de debater "Aspectos centrais da regulação da Inteligência Artificial", a fim de tratar de abordagem principiológica, regulação baseada em riscos, regime de responsabilidade, governança multisetorial, estatuto de direitos, decisões automatizadas, supervisão humana, ética, privacidade e proteção de dados, em cumprimento ao Requerimento nº 4, de 2023, de nossa autoria.

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

Encontram-se presentes no plenário da Comissão: Estela Aranha – está chegando –, Assessora Especial de Direitos Digitais do Ministério da Justiça e Segurança Pública; Dr. Rodrigo Badaró, Conselheiro do Conselho Nacional do Ministério Público; Ana Carla Bliacheriene, Professora da Universidade de São Paulo (USP); Marcela Mattiuzzo, Conselheira do Instituto Brasileiro de Estudos de Concorrência, Consumo e Comércio Internacional (Ibrac); Adriana Rollo, Líder da Comissão Especial de Regulação de Inteligência Artificial da Associação Internacional de Inteligência Artificial (I2AI); Fernanda Rodrigues, Coordenadora de Pesquisa do Instituto de Referência em Internet e Sociedade (Iris).

Encontram-se também presentes, por meio de sistema de videoconferência: Leonardo Netto Parentoni, Professor da Universidade Federal de Minas Gerais (UFMG); Carlos Affonso de Souza, Consultor da Associação Brasileira de Internet (Abranet); Patrícia Peck, Presidente do Instituto iStart de ética e cidadania digital; Cynthia Picolo, Diretora-Presidente do Laboratório de Políticas Públicas e Internet (Lapin).

Agradeço a presença de todos e iremos, daqui a pouco, começar as nossas exposições. (*Pausa.*)

Neste momento, convido todos os nossos expositores para comporem a mesa aqui, diante das suas identificações, por favor. (*Pausa.*)

Nós iremos dar o prazo de até dez minutos para todos os expositores. (*Pausa.*)

Estamos também aguardando, para os próximos instantes, o Presidente desta sessão, que é Vice-Presidente desta Comissão, o Senador Astronauta Marcos Pontes, mas eu gostaria de iniciar essa apresentação, se todos estiverem de acordo. (*Pausa.*)

Possivelmente...

Eu gostaria até que...

Talvez seja melhor fazer daqui, por causa do PowerPoint ou não? (*Pausa.*)

Então, podemos fazer aqui.

Por ordem, se a Dra. Marcela Mattiuzzo já tiver a sua exposição na forma, já pode iniciar pelo prazo de até dez minutos.



SENADO FEDERAL
Secretaria-Geral da Mesa

Com a palavra, V. Sa.

A SRA. MARCELA MATTIUZZO (Para expor.) – Bom, primeiramente, gostaria de agradecer enormemente ao Senador Eduardo Gomes pelo convite, pela oportunidade de estar fazendo essa fala em nome do Ibrac, de contribuir para esses debates.

O Ibrac é o Instituto Brasileiro de Estudos de Concorrência, Consumo e Comércio Internacional. A gente existe desde dezembro de 1992. E o objetivo específico do Ibrac é promover a realização de pesquisas, estudos e debates, exatamente como este, sobre temas relacionados à defesa da concorrência, a comércio internacional e consumo, especialmente. Então, é uma satisfação enorme estar aqui hoje.

Esse movimento, no Senado especificamente e no Congresso Nacional de maneira mais ampla, de regulamentação do tema parece ao Ibrac extremamente relevante. O que a gente gostaria, inicialmente, de deixar claro também é que a gente entende que seria muito importante que uma discussão jurídico-regulatória acontecendo aqui, no Senado Federal e no Congresso Nacional, não deixasse de lado uma discussão estratégica de posicionamento do Brasil sobre o tema da inteligência artificial. Então, a gente sabe que existe a Estratégia Brasileira de Inteligência Artificial, mas a gente também sabe que esse é um documento que precisa ainda de muita concretização do ponto de vista efetivamente estratégico de como o Brasil vai se colocar nesse campo, especialmente no que diz respeito à política industrial e ao desenvolvimento tecnológico propriamente dito. E acho que é importante destacar que muitos dos países que são referência no debate regulatório tiveram essa discussão estratégica ou ocorrendo paralelamente, ou até tendo acontecido antes mesmo de a discussão regulatória começar.

No Canadá, por exemplo, a gente teve uma discussão estratégica começando ali em 2017 ainda; na Coreia do Sul, isso aconteceu em 2019; nos Estados Unidos, que é um exemplo sempre citado, tem um relatório que foi lançado em 2021, antes de o debate, digamos assim, jurídico-normativo efetivamente começar; e, mesmo na União Europeia, que é mencionada geralmente como um parâmetro para o debate regulatório no Brasil – a gente vai ver isso certamente, isso está bem explícito no relatório da Comissão de Juristas –, existe um plano estratégico coordenado para a IA desde 2018, enfim, só ressaltando esse ponto inicialmente.

Falando propriamente, então, sobre os aspectos centrais da regulação da inteligência artificial – acho que o Senador até comentou aqui na abertura da sessão –, o ponto inicial de onde a gente parte é: qual seria a abordagem regulatória cabível para o Brasil neste momento? E, fundamentalmente, acho que a gente pode resumir essa discussão em dois polos: um que entende que, exatamente pela ausência de uma estratégia mais bem articulada e delineada, talvez o melhor caminho agora fosse uma regulação primordialmente principiológica, que elenque objetivos, que elenque caminhos para serem trilhados – por exemplo, centralidade do ser humano na regulação, etc. –, mas ainda não defina regras específicas; e, de um outro lado, a gente tem um entendimento contrário, o de que a criação dessas regras específicas é importante, porque é ela que vai balizar o próprio desenvolvimento tecnológico. Então, faz sentido definir direitos para as pessoas afetadas, por exemplo, e obrigações específicas aos agentes desenvolvedores e operadores da inteligência artificial.

E acho que esse movimento pode claramente ser visto no próprio Congresso Nacional. Então, os projetos de orientação principiológica, inicialmente, como o 21, de 2020, foram, de alguma maneira, discutidos pela Comissão de Juristas ao longo de todo o ano passado e resultaram no 2.338, que é o projeto que fundamentalmente se discute hoje, que efetivamente vai no caminho muito mais de um marco regulatório, com criação de direitos e definição de obrigações.



SENADO FEDERAL
Secretaria-Geral da Mesa

E o que eu gostaria de ressaltar é que o Ibrac entende que é particularmente importante que, se o caminho escolhido pelo Senado for esse segundo caminho, o de um marco regulatório com definição de direitos e tudo mais, esse marco precisa estar muito bem alinhado com os objetivos amplos que o Brasil pretende estrategicamente perseguir. E, obviamente, isso é desafiador na medida em que essa discussão estratégica ainda é algo incipiente.

De toda forma, queria também trazer aqui alguns aspectos da proposta de marco regulatório que a gente acha que merecem maior reflexão.

O PL 2.338 se baseia numa premissa de regulação por risco, então não é uma regulação totalmente isonômica; ela varia de acordo com o risco concreto que um sistema específico representa. Então, naturalmente, as obrigações para os sistemas variam de acordo com o que cada sistema faz ou foi projetado para fazer. Essa é uma lógica subjacente também em discussões em outros países, acho que especialmente no EU AI Act, na União Europeia, e entendemos que ela faz todo o sentido, especialmente pensando que a gente está tratando aqui de uma regulação transversal, aplicável a diversos setores, tipos distintos de sistema, com aplicações totalmente diversas.

O que a gente entende, porém, que precisaria ser aprimorado nesse contexto é o encaixe dessa premissa com a lógica de criação de direitos que está no PL 2.338. O projeto de lei inova; ele faz algo que acho que não existe em nenhum outro país, que é criar direitos específicos para as pessoas afetadas, o que, em tese, se justifica muito por conta dos riscos específicos que o uso da IA pode ensejar. Então, toda essa criação de direitos está, em boa medida, lastreada na lógica do risco de discriminação.

E a gente sabe que todo sistema algorítmico – os de IA não são exceção, pelo contrário – funciona por meio do reconhecimento de padrões. Então, como um sistema de inteligência artificial, como, por exemplo, o ChatGPT, sabe dizer se um idioma é o português ou o alemão? Ele foi apresentado a uma base de dados que contém exemplos do que é a língua portuguesa e do que é a língua alemã e ele foi treinado – geralmente, por um ser humano – para identificar cada um desses exemplos de forma adequada. E ele usa esse conhecimento que ele adquiriu desses padrões para encontrar a resposta de forma comparativa. Isso já gerou diversos problemas.

Se a gente olha, por exemplo, para o caso do reconhecimento facial, como boa parte dos sistemas foi treinado em bases de dados contendo, fundamentalmente, imagens de homens, brancos, ocidentais, os sistemas acabaram tendo dificuldade de replicar essa lógica de identificação, por exemplo, para pessoas negras, levando à não identificação da pessoa como uma pessoa ou à identificação de uma pessoa como sendo outra pessoa e coisas do gênero.

O ponto principal que acho que a IA revela de mais evidente aqui é que existem padrões que a gente claramente pode observar – o caso do reconhecimento facial é um deles – e outros que não necessariamente. Então, a gente já sabe, intuitivamente, que a questão étnico-racial é uma questão que pode levar a uma discriminação problemática do ponto de vista jurídico. Mas e se, ao invés de discriminar com base nesses fatores conhecidos, um algoritmo se depara com um padrão de comportamento que pode parecer para a gente aleatório, mas que o sistema consegue correlacionar com alguma variável? É o que a gente costuma chamar de correlação espúria.

Tem um exemplo famoso na literatura que é um gráfico que demonstra uma curva de dois eventos. O primeiro deles é o número de pessoas que se afogaram ao caírem em uma piscina do ano de 1999 ao ano de 2009, e o outro são os filmes em que o ator Nicolas Cage apareceu no mesmo período de 1999 a 2009. E as curvas são praticamente idênticas. Elas se comportam de maneira muito parecida. E aí a gente olha e fala: bom, o que tem a ver as pessoas que caíram na piscina e se afogaram e o Nicolas Cage? É isso. A gente sabe que isso não tem nada a ver e que isso é simplesmente uma coincidência –



SENADO FEDERAL
Secretaria-Geral da Mesa

o que a gente chamaria em estatística de uma correlação espúria –, mas não necessariamente o sistema vai saber disso. Ele precisa ser treinado a identificar quando isso é uma informação relevante e quando isso não é uma informação relevante e, enfim, assim evitar maiores problemas.

A depender do uso que o sistema faz desse tipo de correlação, a gente pode ter situações discriminatórias – no sentido juridicamente relevante da palavra, de exclusão de direitos, por exemplo. E aí surgiu todo um campo de preocupação com essa temática e de reflexão sobre como lidar com esse problema.

(Soa a campainha.)

A SRA. MARCELA MATTIUZO – O que o Ibrac queria ressaltar aqui – só para finalizar a minha fala – é que a regulação de risco colocada dentro da lógica do 2.338 precisa também conversar muito bem com essa questão dos direitos das pessoas afetadas. Hoje a gente vê uma desconexão entre essas duas coisas. Todo sistema, de qualquer nível de risco, ensaja os mesmos direitos para as pessoas afetadas, e não tem nenhuma orientação no projeto de como isso vai ser adequado aos diferentes graus de risco. E o Ibrac entende que isso é importante de ser endereçado por esta Comissão, pelo Senado e pelo Congresso Nacional.

Muitíssimo obrigada, Senador.

Ficamos à disposição para continuar os debates.

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO) – Muito obrigado, Marcela Mattiuzo, do Ibrac.

Eu quero, neste momento, trazer rapidamente aqui as informações do e-Cidadania, porque, durante as exposições, qualquer um dos nossos convidados pode comentar ou responder eventualmente.

Então, aqui nós temos a pergunta do Richardson A. Silva, de Minas Gerais: "Quais são os princípios fundamentais que devem guiar a regulação da inteligência artificial?".

Temos também a pergunta do Emanuel Polari, da Paraíba, que pergunta: "Qual é o equilíbrio ideal entre o avanço da inteligência artificial e a preservação dos valores éticos da privacidade e da responsabilidade?".

Sandro Júnior, do Rio Grande do Sul: "Quais serão os órgãos responsáveis por acompanharem o desenvolvimento da inteligência artificial?".

Fábio Lima, de São Paulo, pergunta: "Como vamos garantir a integridade econômica do país caso a [...] [inteligência artificial] substitua empregos de baixa renda?".

E o Johan Karl, do Rio Grande do Sul, pergunta: "A regulamentação deve ser apenas no sentido de filtrar conteúdo falso, ofensivo e abusivo. Não deve impedir [também] a inovação". Isso não é um risco?

Também faço aqui um registro importante: estamos todos nós ao vivo na TV Senado e recebemos aqui hoje a visita, nesta Comissão Especial, do Deputado Federal Filipe Martins, que é lá do meu estado, para o nosso orgulho, e foi Vereador na capital, assim como eu. Ele está aqui acompanhando esta importante sessão, já que esse tema também palpita lá na Câmara dos Deputados.

Já temos aqui a Dra. Estela Aranha, que está presente, mas eu vou passar a palavra, de maneira remota, agora, para a nossa querida amiga Patrícia Peck, Presidente do Instituto iStart de ética e cidadania digital, para um tempo de até dez minutos. Em seguida, retomamos aqui a nossa exposição com a Dra. Estela, que já está atrasada para o próximo compromisso. *(Risos.)*



SENADO FEDERAL
Secretaria-Geral da Mesa

A SRA. PATRÍCIA PECK (*Por videoconferência.*) – Estimado Senador Eduardo Gomes, se quiser dar a preferência para a Dra. Estela Aranha, eu aguardo, já que ela tem um outro compromisso, não tem problema nenhum.

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO) – Ela disse que V. Sa. já pode falar, tranquilo.

A SRA. PATRÍCIA PECK (Para expor. *Por videoconferência.*) – Está ótimo.

Então, primeiramente, eu gostaria de agradecer a oportunidade de estar aqui na audiência pública, a convite da Comissão do Senado, para tratar de um tema tão relevante. Eu acredito que o Brasil tem em suas mãos uma oportunidade para criar o marco legal de inteligência artificial. E, como já tem sido feito, já tivemos outras agendas, a reflexão que eu gostaria de trazer aqui aos Congressistas, aos Senadores é justamente que o assunto é complexo. Toda vez que nos deparamos, nesses últimos anos, em legislar sobre inovação tecnológica, sabemos que sofremos o efeito da obsolescência. A tecnologia desafia a Casa Legislativa porque ela tem uma mudança muito mais rápida, assim como a sociedade, do ponto de vista de usos e costumes, do que nós temos de criar as novas leis.

Como primeira colocação, acho que, como uma premissa aqui para reflexão da sociedade, uma coisa é certa: nós já estamos não só utilizando aplicações de inteligência artificial, mas sofrendo efeitos dessas aplicações. Então, uma lei, uma regra que possa abalizar a utilização da inteligência artificial no Brasil já é necessária. Então, se devemos ou não ter uma legislação, a resposta para isso é "sim", devemos, ela se faz necessária.

A segunda questão que eu gostaria de endereçar é: qual a abrangência? Até onde deveríamos alcançar esse marco legal do ponto de vista de ser uma legislação tão técnica e que exige uma participação muito próxima tanto da sociedade civil, quanto dos setores econômicos, principalmente pela característica da tecnologia a que estamos nos referindo? A inteligência artificial por si só não é simples de a gente analisar conceitualmente. Você tem classificações dos tipos de inteligência artificial, dos níveis de autonomia e de como cada aplicação é utilizada em cada um dos setores econômicos. O que isso significa? E eu gostaria de colocar aqui para o Senado que a recomendação seria que é essencial estar no marco legal da IA. Então, o Brasil pode se tornar uma grande referência, trazendo, primeiro, nesse *framework* legal, o que seria essencial.

Daí a gente já pode endereçar uma das perguntas colocadas por um dos participantes que estão ouvindo esta audiência, que é a questão da definição dos princípios. Se existe algo extremamente relevante na relação homem-máquina, no mundo atual, é o princípio da transparência. Nós não podemos ter uma utilização de uma inteligência artificial onde, por dever e obrigação de fábrica, o desenvolvedor não tenha que gerar elementos de identificação e implemento algoritmo para que aquelas IAs tenham que se identificar como sendo um ente robótico ou artificial. E eu digo isso, porque isso pode significar ter que colocar elementos visuais que determinem essa identificação, elementos sonoros, pensando em acessibilidade.

Hoje, conforme avança inclusive o uso de recursos de *deepfake*, inteligência artificial, independentemente de a gente poder ter toda uma satisfação de clima econômico, social, de entretenimento, nós sabemos que já vivemos golpes e fraudes utilizando inteligência artificial. Então, talvez, o combate ao crime nos traga uma audiência maior em uma parte do regramento sobre inteligência artificial ou até mesmo o endereçamento de outras questões que são tão pertinentes, como já colocadas aqui, até porque estão também tratadas em outras legislações. O princípio de tratar dados, desde que não seja de forma discriminatória, já está endereçado na legislação de proteção de dados pessoais, por exemplo. Partindo da premissa de que a inteligência artificial precisa aprender um grande



SENADO FEDERAL
Secretaria-Geral da Mesa

volume de dados, tudo aquilo que a gente pode aproveitar da própria legislação de proteção de dados pessoais já ajuda muito no processo orientativo dessa legislação.

Daí eu gostaria de recortar que muitos dos países – mais de 21 países – que já regularam algo a respeito da inteligência artificial começaram por algum recorte, porque trabalhar um marco que abranja tudo sobre a inteligência artificial é um desafio muito grande. Então, quando se trabalha dentro do que seria criar uma prioridade dentro de alguns setores econômicos e você já tem aplicações de carro inteligente, aplicações de algoritmos na saúde, aplicações de automação para a cidade inteligente, para uso de videovigilância com reconhecimento facial, quando se trabalham já esses eixos temáticos, você facilita conseguir, inclusive, fazer uma avaliação de risco mais assertiva do que a redação mais genérica, que depois tende a não ser tão aplicável ou ainda acabar gerando uma barreira para o desenvolvimento econômico e para a própria inovação.

Aí eu iria colocar aqui, no meu encaminhamento, uma outra questão. Atualmente, para combater o uso da inteligência artificial que finja ser uma pessoa, eu tenho me envolvido em alguns testes. Nessa parte em específico, se observa, por exemplo, que a gente vai precisar usar o máximo de dados para diferenciar uma IA de um ser humano, talvez até a captura de um batimento cardíaco, porque hoje a inteligência artificial não está copiando o nosso coração batendo, mas isso significa que algumas questões tratadas no PL 2.338, como sempre exigir tratar o mínimo de dados, podem não ser aplicáveis quando a gente tem que proteger o ser humano de uma inteligência artificial que já está sendo utilizada com desvio ou com algum nível de alucinação. Então, a gente só tem que ter um cuidado, para que as regras não se tornem depois algo que possa prejudicar a própria defesa do indivíduo, que é um outro princípio, de que a inteligência artificial não possa fazer mal aos seres humanos.

Então, por isso que – e concordo com a fala que foi feita pela colega anteriormente, e aqui representando o Instituto iStart de Ética e Cidadania Digital –, se a redação do PL 2.338, alcançar uma padronização de conceitos, alcançar uma padronização de princípios seguindo as orientações da Unesco, da OCDE, se trouxer ali um compromisso de regras que já deveriam vir de fábrica... Por exemplo, se não pode o ser humano alegar desconhecimento da lei para evitar uma responsabilidade, também não poderia uma inteligência artificial, um algoritmo alegar desconhecimento da lei. Eu não poderia permitir que um carro inteligente saísse de fábrica sem conhecer o Código de Trânsito nacional. Então, o dever de conhecer as leis vigentes para que aquela IA possa se relacionar, o dever de denúncia, de reportar e colaborar com autoridade, nós temos deveres que a gente já poderia prever como obrigações de fábrica para qualquer inteligência artificial.

É assim que eu encerro, e agradeço a oportunidade de estar aqui na audiência. Muito obrigada.

O SR. PRESIDENTE (Eduardo Gomes. Bloco Parlamentar Vanguarda/PL - TO) – Muito obrigado. Quero agradecer à nossa querida amiga Patrícia Peck.

Passo, neste momento, a palavra à Dra. Estela Aranha, para sua participação.

O Senador Marcos Pontes deve estar chegando à Comissão, nós estamos revezando. Quero também reforçar aqui o pedido aos meios de comunicação da Casa, TV Senado, Rádio Senado: estamos pedindo muita colaboração deles na versão, na construção da informação objetiva de todas as audiências públicas, afinal de contas, foi uma decisão do Presidente Carlos Viana e de todos os membros desta Comissão que nós fizéssemos todas as audiências necessárias possíveis para reforçar o posicionamento de comunicação, de abertura a todos que queiram participar desse debate, tendo em vista a necessidade de legislação.

Eu tenho brincado muito com amigos nossos: a gente vai ter que fazer uma legislação precisa, minimalista e viva, porque é um assunto que, entre uma Comissão e outra, pode mudar completamente.



SENADO FEDERAL
Secretaria-Geral da Mesa

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

A SRA. ESTELA ARANHA (Para expor.) – Bom dia, Senador. Queria muito agradecer o convite desta Comissão e, em especial, elogiar o trabalho desta Comissão, das audiências públicas. É muito importante a gente debater este tema para aprofundar o debate. Agradeço também à Adriana, que está tornando isso possível, da assessoria da Comissão, e a todos os meus colegas aqui de mesa também.

Justamente mudou minha data aqui de apresentação porque eu estou voltando do Chile, justamente de um evento latino-americano e do Caribe sobre inteligência artificial. Esse é um debate do mundo inteiro, a gente tem diversos órgãos internacionais que estão debatendo a questão de inteligência artificial. Certamente nós teremos padrões e parâmetros internacionais saindo de alguma forma, porque isso é um desejo de todos os países. Nós estávamos fazendo lá uma discussão de ética junto com a comunidade latino-americana e do Caribe, mas hoje, por exemplo, a ONU, por volta de 1h30 aqui, no Brasil, vai lançar um comitê de aconselhamento ao Secretário-Geral sobre esse tema, justamente para regras de governança global, eventualmente até um órgão de governança global.

Também o Reino Unido agora, no final do mês, está fazendo também um evento do que eles chamam de *frontier*, que é a IA de fronteira, aquela que pode pôr em risco segurança pública, defesa nacional, entre outras coisas. Então, também o G7 está fazendo discussões lá no processo de Hiroshima e vai tentar trazer nesta gestão algum documento em relação à IA generativa e a esses riscos mais de fronteira – como tem todo um debate de acordo da União Europeia e Estados Unidos. Então é só para ressaltar que não é uma especificidade do Brasil: o mundo inteiro está discutindo isso e discutindo a necessidade de regulação em diferentes níveis, tanto as questões relacionadas à questão que é de fronteira, que resvala na segurança pública, resvala nesses riscos mais graves, quanto também a outros riscos que têm a ver com riscos em relação a direitos humanos, dos quais eu vou falar um pouco aqui.

Então é a importância da regulação, tendo em vista que a primeira onda de digitalização, de modo geral no mundo, teve uma postura de abstenção em relação a uma regulação mais forte. E hoje a gente tem essa percepção, do mundo inteiro, da necessidade de você ter alguns parâmetros regulatórios de modo geral, para que não aconteça como nessa primeira onda de digitalização.

A outra questão que a gente... Enfim, começando aqui um pouco no ponto do debate, a gente está falando... Primeiro, quando a gente está falando de princípios – e aí a questão é em relação principiológica –, todos concordam de alguma forma, obviamente, com os princípios: a tríade constituinte do nosso Estado democrático de direito, como o da maioria das democracias liberais; a produção de direitos humanos; a preservação das liberdades fundamentais; a produção da democracia; a redução da desigualdade; a não discriminação; a transparência; a responsabilidade; a centralidade no elemento humano; a mitigação de risco; a autonomia humana, ao que a gente chama de *trustworthy*, a confiança, no sentido técnico, de a inteligência artificial entregar o que promete e com segurança. A dificuldade nossa é fazer com que esses princípios gerais justamente tragam esses direitos, essas garantias fundamentais, esses valores que devem nortear o desenvolvimento da tecnologia, e não só esses princípios como princípios legais ou como princípios de *design*. É isto que a gente tem que tentar fazer aqui: como é que a gente vai criar esses mecanismos capazes de operacionalizar todas essas regras estatutárias, principiológicas, mas para que elas sejam parâmetros obrigatórios para o processo de desenvolvimento e dos usos também da inteligência artificial; como é que essas regras vão ser adequadas a um contexto que requer um sistema de normas, em algumas instâncias granular e bastante dinâmico, como o próprio Senador falou agora, voltado também para diferentes setores, tanto da economia quanto da tecnologia, que têm sua própria linguagem, que têm valores próprios, que têm parâmetros próprios e essas questões setoriais.



SENADO FEDERAL
Secretaria-Geral da Mesa

E isso é mais complexo ainda, pois a gente está falando de trazer – a gente sempre traz este debate dos valores internacionais porque a gente sabe que não existem padrões únicos –, inclusive, trazer esses princípios. Os próprios tribunais interpretam todos esses princípios de uma forma bastante diferente no mundo inteiro. A gente tem um exemplo, o da liberdade de expressão, que nos Estados Unidos é interpretada de uma forma; na Europa, de outra; no Brasil, nós temos outra interpretação. Então, como é que a gente vai fazer isso, não é? Esses desafios é que nós precisamos encontrar nessa regulação, encontrar uma ferramenta de governança, como fazer um *enforcement* dessas regras também e, ao mesmo tempo, ter essas normas que têm que garantir mesmo esses direitos fundamentais, e instrumentos de *soft law*, que a gente chama, de autorregulação, entre outras coisas, que também são necessários para se somarem a todo esse ecossistema. Então, acho que a gente podia pensar num ecossistema regulatório, que tem uma regulação mais geral, transversal, que dá regras gerais e de garantias desses valores fundamentais e de como a gente vai trabalhar isso dentro do sistema de governança para ter toda essa granularidade, essa flexibilização e trabalhar esses outros elementos.

A Comissão de Juristas, como a colega falou, trouxe uma novidade. E eu acho que é uma novidade muito importante a que a Marcela trouxe, que é a atenção que nós demos aos padrões substantivos prometidos nesses princípios, nas garantias processuais que articulam as regras e procedimentos formais por meio dos quais esses direitos substantivos são criados, exercidos e aplicados nessa abordagem que a Comissão tomou, que é mista e que se chama *risk-based approach*, uma abordagem baseada nos riscos, e na *rights-based approach*, que é uma abordagem baseada nos direitos, dentro do debate internacional.

Isso aconteceu primeiro também porque a gente não tem um ecossistema de alguns direitos garantidos, como, por exemplo, a União Europeia, que não traz direitos no EU AI Act, mas o próprio GDPR traz um monte de direitos, como a revisão humana, entre outras coisas, algumas coisas de devido processo, tem o GAC para outras coisas, enfim, tem outros complexos direitos que aqui a gente não tem. Então, como a gente está visando regular não uma tecnologia, mas os efeitos sociotécnicos dessa tecnologia e, principalmente, os impactos deles nos direitos e garantias fundamentais, criar esses novos direitos é muito necessário no sentido dessa proteção substantiva desses direitos fundamentais. É muito importante isso.

Então, eu vou falar um pouco desses direitos, um pouco fazendo essa... Na verdade, são direitos e não têm... Quando são direitos, eles são gerais, não vão ser aplicados em caso, por exemplo, dependendo do risco, no sentido de que é um direito geral e que o impacto tem que ser equivalente àquilo que o uso da tecnologia vai impactar, a responsabilidade que se tem que ter de acordo com o nível de impacto, com o nível de risco, com o efeito legal ou significativo. Eu vou dar um exemplo aqui. A gente está falando de transparência, que é diferente em diferentes sistemas. Se eu tenho um sistema, por exemplo, de trabalho, em que o contrato de trabalho virou algoritmo de regras de demissão praticamente, de você entrar, regras da sua remuneração... Mesmo que não seja um trabalho do regime CLT, é trabalho, é uma forma de trabalho, e a ideia é o trabalho decente, não é? Se essas regras estão sendo postas pelo algoritmo, elas têm que ser, sim, transparentes, e aquele trabalhador tem que saber quanto ele vai receber por determinado trabalho. Obviamente, enfim, tem granularidades e diferenças nesses direitos.

Vou falar um pouco desses direitos. Outro direito protegido é o livre desenvolvimento da personalidade, que está protegido e que deriva da liberdade cognitiva, que é assegurada pelos direitos, neste projeto de lei, associados à informação, à compreensão das decisões tomadas por sistemas de



SENADO FEDERAL
Secretaria-Geral da Mesa

inteligência artificial, como direito à informação prévia quanto às suas interações nos sistemas de inteligência artificial. Outro direito protegido é a garantia substantiva do devido processo legal e também da autonomia humana, que deriva do direito de personalidade do qual se extrai o direito à explicação sobre a decisão automatizada, recomendação e previsão de tomada de decisão nesses sistemas de inteligência artificial, contestação de decisão ou previsões...

(*Soa a campainha.*)

A SRA. ESTELA ARANHA – ... da inteligência artificial que produzam efeitos jurídicos, entre outras coisas.

Só finalizando aqui por causa do tempo, trago também o direito à igualdade, à não discriminação, no qual se extraem esse direito à não discriminação algorítmica que está no projeto de lei e a correção de vieses discriminatórios diretos, indiretos, ilegais, abusivos, incluindo aí aquilo que a gente chama de *disparate impact*, o que também a colega Marcela já explicou muito bem. Ele é considerado um artefato próprio do processo tecnológico do *machine learning*, do aprendizado de máquina. Então, é necessário que a gente tenha esse direito garantido, porque, se nada for feito, na verdade, a gente tem a tendência a ter essas *bias*.

E, por fim, um direito que é fundamental é o direito à inferência razoável, para proteger de todas as decisões, recomendações ou previsões que sejam discriminatórias, irrazoáveis, que atuam contra a boa-fé, que sejam fundadas em dados abusivos, inadequados, baseados em métodos imprecisos, porque *machine learning* nada mais é quando a gente está falando de decisões automatizadas nesse aspecto de probabilidade...

(*Soa a campainha.*)

A SRA. ESTELA ARANHA – Então, a gente precisa ter métodos estatísticos confiáveis. E que também não considere adequada a individualidade com as características desses indivíduos.

Eu queria muito ressaltar a necessidade desse enfoque aos direitos, em especial esses direitos de que a gente tratou aqui.

Obrigada.

E bom dia, Senador.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Bom dia, bom dia a todos.

Eu queria cumprimentar todos os presentes aqui, todos os nossos debatedores, apresentadores, todos aqueles que nos assistem pela TV Senado também e pelas redes do Senado.

Só vou enfatizar aqui para quem quiser participar, que está nos acompanhando pelas redes, que pode participar através do Portal e-Cidadania, www.senado.leg.br/ecidadania, ou pelo telefone 0800 0612211. Participem, tragam suas questões. Já temos algumas aqui.

Eu gostaria especialmente de agradecer a todos os nossos apresentadores.

Eu estou com uma sequência aqui e, só para colocar... Quanto tempo está sendo para cada um? Desculpe... (*Pausa.*)

São dez minutos, dez minutos para cada um. Então, pela sequência que o Eduardo me deixou, eu chamo agora o Professor da Universidade Federal de Minas Gerais... (*Pausa.*)

É a Fernanda primeiro aqui? Então, desculpe, Leonardo. Desculpe, acionei errado o motor aí.



SENADO FEDERAL
Secretaria-Geral da Mesa

Para a nossa próxima apresentação, eu gostaria de convidar a Sra. Fernanda Rodrigues, Pesquisadora e Coordenadora de Pesquisa do Instituto de Referência em Internet e Sociedade. A senhora tem dez minutos.

Aqui é mais fácil de controlar o tempo, porque aparece o tempo aqui. Para quem estiver *online*, eu peço para que controlem o tempo, porque não vão aparecer as referências. Obrigado, por favor.

A senhora tem a palavra.

A SRA. FERNANDA RODRIGUES (Para expor.) – Bom dia a todos, todas e "todes".

Primeiramente, eu gostaria de iniciar a minha fala agradecendo pela oportunidade de estar aqui, principalmente na pessoa do Senador Marcos Pontes, que está aqui conosco, mas também na pessoa do Senador Alessandro Vieira e de sua assessoria, que pôde indicar o Iris para participar aqui deste espaço com vocês.

E, para quem não conhece, para mim é uma honra estar aqui representando o Instituto de Referência em Internet e Sociedade (Iris). Ele é um centro de pesquisa interdisciplinar e independente que foi fundado em 2015 e que tem como missão principal fortalecer os direitos humanos no espaço digital, principalmente através da comunicação, incidência política e pesquisa. A gente desenvolve projetos de pesquisa não só na área da inteligência artificial, que é o tema que a gente vai discutir aqui hoje, mas também em regulação de plataformas, criptografia, proteção de dados, inclusão digital, e a gente espera hoje trazer uma contribuição muito importante aqui para esta Casa.

A gente parabeniza o trabalho da Comissão, porque, desde que a gente tem acompanhado, enquanto Iris, principalmente a partir da Estratégia Brasileira de Inteligência Artificial, essas discussões, e a gente pôde participar inclusive das audiências públicas na Câmara dos Deputados sobre o PL 21, de 2020, e também, no ano passado, no Senado Federal, a gente não desconhece o quão difícil é a tarefa que esta Casa está empreendendo no sentido de regular uma tecnologia que não somente está em constante evolução, como também está em uma evolução muito acelerada. Só que a gente acredita que isso não é motivo para se furtar à regulação desse objeto.

A gente acredita que é falsa essa dicotomia entre inovação e regulação, na medida em que estabelecer esses critérios e parâmetros para a atuação dos agentes envolvidos, bem como os direitos e garantias para pessoas afetadas e para cidadãos e cidadãs no geral, pode fornecer segurança jurídica e permitir que essa inovação tecnológica se dê de maneira ética e responsável.

Nesse sentido, a gente acredita, no entanto, que apenas uma regulação por princípios e valores não é suficiente para a gente endereçar a série de questões relacionadas à IA.

Infelizmente, a gente tem notícia diariamente de violações e prejuízos que são causados por sistemas da inteligência artificial. E aqui eu posso citar como exemplo notícias falsas, discriminação algorítmica e muitas outras situações. Mas o que a gente gostaria de destacar é que, para além de erros e falhas, a gente está falando aqui de pessoas, muito provavelmente de pessoas que têm a sua esfera de direitos pessoais afetada por esse tipo de tecnologia.

Até gostaria de compartilhar, porque eu acho que é uma provocação que norteia a minha fala, principalmente porque, além de coordenadora de pesquisa e pesquisadora do Iris, eu sou doutoranda em Direito na Universidade Federal de Minas Gerais e eu estou fazendo uma disciplina esse semestre que é sobre racialização de existências, junto com o Prof. André Luiz Freitas... Ele fala muito sobre... A gente discute muito nessa aula sobre quais existências importam para o direito. Eu acho que aqui a provocação é: quais existências importam e precisam constar na regulação da inteligência artificial para que a gente tenha uma regulação ética e responsável?



SENADO FEDERAL
Secretaria-Geral da Mesa

Nesse sentido, a gente destaca a importância do Projeto de Lei 2.338, que é fruto do trabalho da Comissão de Juristas, na medida em que ele propõe uma regulação baseada em riscos. O que a gente quer apontar, porém, é que essa matriz de risco deve ser pensada principalmente considerando o contexto e as particularidades do Brasil. É preciso que a gente supere essa ideia da tecnologia como um fim em si mesma e a gente busque compreender o que pode ser automatizado, o que faz sentido ser automatizado e o que não pode ser automatizado.

É claro que a inovação tecnológica é absolutamente importante para o avanço e desenvolvimento da economia do país, mas é necessário que a gente pense quais tecnologias fazem sentido no nosso contexto.

E aqui no Brasil, eu gostaria de destacar, principalmente... E aí, a gente tem que fazer uma seleção do que falar aqui, por causa do tempo. Mas, no cenário brasileiro, o que eu gostaria de destacar – que muitos já conhecem, mas eu acho que o óbvio, às vezes, é importante que seja dito – é que existe um público muito específico, que tem sido afetado por tecnologias, mas também pelo preconceito na sociedade, que é a população negra. E é importante a gente denotar essa questão racial dentro da inteligência artificial não somente porque o Brasil possui um histórico escravocrata conhecido, mas porque a gente tem diversos estudos reforçando o quanto as tecnologias de inteligência artificial especificamente podem ser especialmente nocivas para a população negra. E aí, ter isso como norte implica algumas considerações.

A primeira delas é que existem tecnologias que não podem ser aceitas no nosso contexto, porque os seus prejuízos são muito superiores aos seus eventuais benefícios. O reconhecimento facial em espaços públicos, principalmente para fins de segurança pública, e também o policiamento preditivo são dois que a gente gostaria de destacar aqui. Principalmente em relação ao reconhecimento facial, a gente tem uma campanha nacional chamada Tire meu Rosto da sua Mira, que pauta o banimento dessa tecnologia para fins de segurança pública, mas, para além disso, a gente tem farta literatura demonstrando o quanto a segurança pública e o sistema penal brasileiro como um todo são baseados estruturalmente no racismo. Isso significa que qualquer tecnologia ou qualquer medida, estratégia voltada para essas áreas específicas precisa considerar quais são as principais pessoas afetadas nesse sentido.

Eu gostaria aqui de destacar e falar melhor sobre outras tecnologias, mas, como o tempo está correndo, eu gostaria só de recomendar a leitura da nota técnica da Coalizão Direitos na Rede, da qual o Iris faz parte, que foi elaborada pelo Laboratório de Políticas Públicas e que destaca, além dessas duas tecnologias inaceitáveis, sistemas de score de crédito, armas letais autônomas e sistemas de reconhecimento de emoções, que merecem igual vedação.

Com relação aos demais sistemas – aí há outro ponto que a gente gostaria de destacar sobre a regulação da IA no Brasil –, é fundamental que a regulação preveja robustos mecanismos de fiscalização e validação a fim de garantir que todo o ciclo de vida da tecnologia seja pautado por valores antidiscriminatórios, principalmente antirracistas, ainda mais considerando que, no nosso país, mais da metade da população se autodeclara preta ou parda; então, a gente está falando da maioria, em números pelo menos. Então, é importante que a gente tenha mecanismos de validação e fiscalização, e aqui a gente destaca a avaliação do impacto algorítmico e auditorias de sistemas algorítmicos, mas que eles sejam concebidos não somente como mero *checklist* para empresas, mas que sejam instrumentos efetivamente capazes de avaliar a fundo os possíveis riscos e impactos que podem advir da IA antes de entrar no mercado e também, após a sua implementação, um acompanhamento e monitoramento para ver se algum prejuízo ou dano é causado para a sociedade.



SENADO FEDERAL
Secretaria-Geral da Mesa

Em relação ao racismo, a gente gostaria de destacar inclusive um texto recente publicado pelo pesquisador e referência nacional Tarcízio Silva – que também vai poder estar aqui nas audiências da semana que vem, se não me engano –, que ele publicou recentemente no *Portal Jota Info*, em que abordou a necessidade de lentes antirracistas e interseccionais para identificar impactos algorítmicos em conformidade com princípios e mecanismos antidiscriminatórios. Ele afirma inclusive ser necessária a ênfase em quatro pontos principais – e aqui eu abro aspas –: "compreensão de sistemas algorítmicos dentro das estruturas de relações sociais; reconhecimento do caráter estrutural do racismo e seus impactos; centralização dos compromissos do Estado brasileiro contra o racismo e outras formas de discriminação; e reconhecimento dos impactos do colonialismo e colonialidade digital".

Por essas razões, a gente entende como primordial que o texto da regulação da IA reconheça expressamente a necessidade dessa postura antirracista por parte das empresas. Inclusive, o Brasil ratificou a Convenção Interamericana contra o Racismo, Discriminação Racial e Formas Correlatas de Intolerância, no início do ano passado, com o *status* de emenda constitucional. Por isso, é fundamental que sejam tomadas medidas proativas para evitar esse preconceito na sociedade brasileira. Apesar de a gente ter realmente regulações específicas que falam sobre discriminação, eu retorno aqui a provocação no sentido de quais existências importam para a regulação da IA, porque, quando a gente leva isso para o texto específico da regulação da inteligência artificial, mesmo sendo algo que já está disposto no ordenamento jurídico brasileiro, a gente dá uma mensagem para a sociedade de que a gente está colocando isso como pauta prioritária para fins do desenvolvimento desses temas de IA no Brasil.

E aqui eu falo que é necessário estar na lei...

(Soa a campainha.)

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Obrigado, Sra. Fernanda Rodrigues, Coordenadora de pesquisa e pesquisadora do Instituto de Referência em Internet e Sociedade.

Na sequência, eu passo a palavra, agora sim, para o Leonardo Netto Parentoni, Professor da Universidade Federal de Minas Gerais.

O senhor tem dez minutos, e eu peço para controlar o tempo, porque só vai aparecer no final uma voz de "15 segundos" – uma voz muito convincente, feminina.

O SR. LEONARDO NETTO PARENTONI (Para expor. *Por videoconferência.*) – Muito obrigado. Eu agradeço a honra de participar desta audiência pública.

Cumprimento o Senador Marcos Pontes, o Senador Carlos Viana, meu conterrâneo de Minas Gerais, todos os expositores e todos que nos assistem.

Minha fala é muito direta, até pelo tempo, e eu reporte um ponto específico: como a IA interfere na decisão humana e quais são os reflexos disso para uma possível regulação?

A fonte científica do que eu falo – e é apenas uma manifestação técnica – é esse artigo que está disponível gratuitamente nesta plataforma, no ResearchGate. Basta pesquisar no ResearchGate o meu nome, não precisa ter *login* e senha. Esse artigo já foi publicado em Portugal e na Itália e até o final deste ano deve ser publicado em Rússia, China, Uruguai e Estados Unidos.

Então, qual é a ideia central desse texto? Há muita coisa que todos nós ouvimos quase que diariamente: o ideal é que a inteligência artificial sempre supere a capacidade humana, sempre seja mais transparente do que o ser humano e, no final, seja sempre mais justa. Isso é dito muitas vezes na literatura técnica ou na imprensa.



SENADO FEDERAL
Secretaria-Geral da Mesa

O meu ponto central é que o erro está no ser. Não tenho dúvidas de que são três objetivos muito importantes, mas é preciso analisar caso a caso. E aqui o meu ponto é contraintuitivo. Eu tive o cuidado de ver o vídeo das audiências públicas anteriores para que eu pudesse trazer um ponto que, ao menos na minha percepção, ainda não havia sido discutido, que é basicamente o seguinte: uma inteligência artificial pode ser boa, ainda que não cumpra algum desses três requisitos. E eu o fundamento numa frase de 1979: "Todos os modelos de IA falham, mesmo assim muitos são úteis". Qual é a ideia central? Existem contextos de fato específicos em que não faz sentido exigir da inteligência artificial uma precisão ou acurácia igual ou muito superior à do ser humano. E as vantagens disso são: reduzir o ciclo de desenvolvimento de produtos ou serviços; reduzir o custo; fomentar a inovação e, consequentemente, gerar desenvolvimento econômico e social. Isso, obviamente, se aplica a alguns contextos, não a todos, e eu gostaria de frisar esses contextos.

Por exemplo, em casos de atividade com risco inerente à integridade ou à vida do ser humano, parece-me fazer sentido que substituamos a atividade humana por uma atividade de um robô ou de um sistema de inteligência artificial. Por hipótese, um robô que desarma bombas e um ser humano altamente especializado em desarmar bombas. Ainda que o ser humano tenha uma taxa de acerto de 90%, parece-me fazer sentido usar um robô com a taxa de acerto de 50% ou 60%, porque, em última análise, quando cairmos nesse percentual menor, em vez de perder uma vida, perderemos uma máquina.

E isso se aplica também para qualquer atividade inerentemente insalubre. Outro exemplo: limpeza de tanques de óleo em navios e em empresas. Isso pode ser feito por uma equipe humana, e normalmente é feito, ou pode ser feito por sistemas de inteligência artificial, ou eventualmente por ambos. Parece-me que, ainda que essa atividade executada com o auxílio de um robô demande mais tempo, não seja tão eficiente, o valor que é preservar a integridade e a vida de um ser humano se justifica.

Então, parece-me um mito exigir-se que a inteligência artificial seja sempre, em qualquer caso, muito superior ao ser humano. Há várias situações de fato em que ela é útil e até deveria ser incentivada, ainda que com precisão menor ou igual à do ser humano. Com precisão igual, nós temos todas as atividades meramente burocráticas. Se uma IA tem uma taxa de acerto semelhante à de um ser humano, melhor usar a IA do que o ser humano, porque liberamos para nós pessoas e tempo livre para outras atividades mais importantes.

Em um terceiro nível, há áreas estratégicas, por exemplo, segurança pública e saúde, em que, nessas sim, eu concordo – e sustento cientificamente – com que a legislação deveria exigir o mais alto grau de acurácia e de precisão disponível, porque, aí sim, estamos lidando com áreas em que ou o trabalho é feito por um ser humano, ou a premissa para se permitir que uma inteligência artificial substitua o trabalho humano nessas áreas é a garantia de que ela é muito mais precisa. Numa cirurgia robótica, ou é comprovado que o robô que realiza aquela cirurgia tem uma taxa de acerto muito superior à de uma equipe médica humana, ou ao menos eu, Leonardo Parentoni, prefiro ser operado pelo bom e velho médico humano. As vantagens desse raciocínio eu comentei com os senhores.

Qual é o critério que me parece que pode ser utilizado neste PL 2.338, ou em qualquer projeto de lei sobre o tema, para saber qual desses três níveis nós vamos utilizar?

Parece-me que o primeiro passo seria dividir o sistema de inteligência artificial no caso concreto em uma das seguintes classificações: ou ele é o sistema intermediário, ou inicial, ou avançado. Do princípio para o fim: inicial, intermediário ou avançado. Cada um desses sistemas tem características de funcionamento diferentes. Alguns deles só acessam o banco de dados estático; outros acessam dados



SENADO FEDERAL
Secretaria-Geral da Mesa

em tempo real. Alguns deles só funcionam por iniciativa de um ser humano; outros funcionam automaticamente a partir do momento em que ligados. Alguns deles só sugerem uma decisão para o ser humano; outros substituem o próprio ser humano na tomada de decisão. São três conceitos e três situações absolutamente diferentes, que eu tentarei mostrar com exemplos neste momento.

Parece-me que a regulação, seja ela brasileira, seja de qualquer país, deveria levar em conta essas diferenças para que efetivamente resolva o problema sem prejudicar a inovação ou o desenvolvimento do mercado. A premissa é: quanto maior a interferência do sistema de inteligência artificial na tomada de decisão humana, maior deve ser o rigor da legislação com aquele tipo de sistema. E, ao contrário, se o sistema não interfere substancialmente na tomada de decisão humana, não há por que a legislação criar requisitos e barreiras para aquele sistema, porque, afinal, a decisão continuará sendo humana.

Nesse nível 1 de automação, por exemplo, nós temos os assistentes vocais de celular. Quando um Exmo. Senador pergunta ao seu telefone "me indique um restaurante por perto", certamente a primeira opção que vai ser indicada não é porque é a melhor comida ou a preferência do Senador, mas é porque aquele restaurante tem um acordo comercial com o fabricante do telefone e paga ao fabricante do telefone para que, em perguntas feitas por usuários próximos, o telefone indique aquele restaurante, e não um concorrente igualmente bom.

O que eu sustento é que esse tipo de acordo comercial não só é lícito, como também pode ser sigiloso, e é hoje sigiloso na prática do mercado. Seu assistente vocal não diz que está recomendando porque aquele restaurante patrocina Apple, Samsung ou quem quer que seja. Por que isso é lícito? Porque, ao final do dia, não há qualquer substituição da decisão humana. É o próprio usuário que decide se vai àquele restaurante, a qualquer outro, ou se mesmo não vai almoçar naquele momento. A decisão final continua sendo 100% humana.

No último nível, no nível máximo de automação, nós temos uma inversão dessa ordem. Numa cirurgia robótica, a equipe médica, dependendo da técnica utilizada, fica apenas de plantão. Ela não toca no paciente após a anestesia. É o próprio robô que desempenha tudo. Evidentemente, neste caso, é indispensável a intervenção do legislador para proibir esse tipo de acordo comercial, porque ele vai guiar uma cirurgia, que pode redundar em vida ou morte, não pelo interesse do paciente, mas pelo interesse comercial do fabricante. E isso precisa ser decidido pela legislação.

Idem para o carro autônomo. Ele deve ter um nível maior de previsão na lei, porque não é o ser humano que toma as decisões, é o veículo. Portanto, para casos de maior risco, exige-se maior transparência, acurácia e explicação, e podem ser proibidos determinados acordos comerciais.

Em sentido contrário, não me parece fazer sentido que a legislação intervenha em sistemas de inteligência artificial que apenas recomendam, mas não substituem a decisão humana, porque intervir nesses casos seria sufocar ainda mais as pequenas e médias empresas brasileiras, que lutam com muita dificuldade para se manter no mercado e que, ao final do dia, o sistema que elas proporcionam apenas recomenda, mas não substitui a decisão humana.

Utilizei 9 minutos e 50 segundos.

Agradeço novamente a oportunidade de participar desta audiência pública e me coloco à disposição, caso haja interesse em fazer qualquer tipo de pergunta.

Muito obrigado.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Obrigado, Prof. Leonardo Netto Parentoni. Não só parabéns pela apresentação, mas parabéns pelo controle do tempo também. Muito bom.



SENADO FEDERAL
Secretaria-Geral da Mesa

Eu gostaria de fazer alguns comentários, mas antes também agradecer aos ouvintes da Rádio Senado, que também nos acompanham ao vivo nesta audiência pública.

Esse é um tema extremamente importante por várias razões: primeiro, porque já está na vida de todos nós de certa forma, a inteligência artificial já participa de muitos dos sistemas que nós usamos com frequência no dia a dia; e também, quando se trata de uma legislação dessa natureza, é importante – só resumindo algumas coisas que a Sra. Fernanda falou também – ter sempre o ser humano no centro dessa legislação. Qualquer tecnologia disruptiva, como é o caso da inteligência artificial, tem aplicações muito positivas. Eu vou usar como exemplo a tecnologia nuclear, que tem aplicações muito positivas para diagnósticos – todos os dias isso é utilizado nos hospitais, etc. –, medicamentos, etc. E qualquer tecnologia disruptiva grande pode também ser utilizada para o lado ruim, como a bomba atômica, por exemplo, que é ruim para todo mundo, e a inteligência artificial não foge desse padrão.

Então, é importante que nós prestemos atenção nisso, mas sempre com esse enfoque que o Leonardo citou, que eu acho importante. Nós não podemos, vamos dizer assim, restringir a tecnologia, mesmo que ela vá acontecer de qualquer forma. Se não acontece no Brasil, vai acontecer em outros países, e a gente vai importar a tecnologia com os problemas inerentes, inclusive da análise de dados, já que a inteligência artificial trabalha com muitos dados, e a interpretação desses dados tem a ver também com o aprendizado que é feito por essa máquina. Então, isso... (*Pausa*)

Está o.k. Obrigado.

Desculpa. Eu perdi um pouquinho o foco aqui.

Ela tem, então, uma série de implicações, e você não consegue, obviamente, contornar ou prever, mesmo porque essa é uma tecnologia em desenvolvimento. Nunca vai ser possível prever, ou seja, criar alguma coisa preditiva sobre como vai ser essa tecnologia, mas nós podemos criar situações ou imaginar situações de uso em que, aí, dentro da correspondência e das necessidades de proteção do ser humano, como centro de tudo isso, a gente pode, sim, colocar os devidos cuidados. Isso envolve a questão da utilização de dados das pessoas, as questões éticas, as questões de discriminação que podem ocorrer – existe o potencial de ocorrer – de acordo com a utilização, o aprendizado dessas máquinas. Então, isso é a parte que nós podemos fazer com certeza, na parte de utilização. Na parte do desenvolvimento, não dá, mesmo porque, se a gente fizesse uma lei dessa forma, no dia seguinte ela teria que ser já atualizada, porque o desenvolvimento é praticamente impossível de se prever e também é impossível de se acompanhar com a lei.

Um outro ponto importante – eu tomei nota aqui também – a ser falado é com relação ao foco positivo. Obviamente, nós temos todo um receio de tudo o que é novo. Isto faz parte do ser humano: tudo que é novo causa medo. E a gente não pode deixar esse medo ofuscar as vantagens da utilização. A inteligência artificial tem uma série de vantagens, e uma vantagem muito grande, na qual, em parceria com a inteligência humana e as emoções humanas, ela pode acelerar muito o desenvolvimento de coisas de que a gente precisa há muito tempo. Há problemas do nosso planeta que a gente precisa resolver. Cada vez mais, a gente fala de mudanças climáticas, a gente fala da necessidade da produção de alimentos, preservação, utilização de água, doenças, epidemias, pandemias; tudo isso ela pode acelerar e muito – a gente já vê as aplicações sendo feitas –, com o cuidado, lembrando também o que o Leonardo falou, das decisões.

A decisão humana é basicamente focada na emoção. Nós tomamos decisões justas através da emoção. A decisão lógica, às vezes, não se encaixa na decisão justa.



SENADO FEDERAL
Secretaria-Geral da Mesa

Então, é um cuidado que nós temos que tomar sempre: de o ser humano ser o decisor, de as coisas importantes serem sempre decididas por nós. A inteligência artificial nos ajuda a congrega um monte de dados, a trazer certas sugestões, mas o ser humano é o responsável, em essência.

Na profissão, por exemplo, de piloto, um piloto comercial, com 200 pessoas a bordo: por mais que exista piloto automático, por mais que exista inteligência artificial para auxiliar dentro da aeronave, quem é o comandante? É o piloto humano. E ele toma as decisões, ele é o responsável em si por aquilo lá.

É importante a gente sempre ter isso aí em mente.

É lógico que existem alguns sistemas, como o Leonardo lembrou, como carros autônomos, etc. O cuidado em se deixar a vida humana às decisões artificiais, vamos dizer, à decisão de uma máquina, tem que ser extremamente grande. E isso tem que ser restrito aos fatos, às aplicações muito específicas, sob um controle muito específico, sob um teste muito específico.

Para ter uma ideia, para você certificar um avião para ele pousar automaticamente, existem sistemas? Existem. Eu sou piloto de teste de aviões, eu já fiz isso aí de pousar o avião automaticamente. Mas, obviamente, você está em teste, você está atento para assumir aquele avião em questão de segundos e tomar a providência adequada, com o seu conhecimento, a sua capacidade, a sua competência técnica para fazer aquilo. E, para que um sistema desse seja certificado para poder voar com passageiros, esse sistema tem que passar por inúmeros testes, demonstrando que ele é preciso, e em números absurdos, como um acidente a cada dez à nona pousos, por exemplo – um um com nove zeros –, ou seja, é muito difícil você certificar uma coisa dessa. E tem razão de ser, obviamente, porque a gente está tratando de vida humana.

Passando agora, na sequência, eu gostaria de aproveitar e anunciar a presença do Senador Izalci Lucas aqui conosco. É um dos grandes defensores da ciência e da tecnologia. Eu, como Ministro de Ciência e Tecnologia, trabalhei muito com o Senador Izalci aqui; e continuo agora, tenho a honra de trabalhar com ele aqui no Senado.

E eu queria aproveitar, Senador Izalci, se puder nos acompanhar aqui também na mesa...

Enquanto o Senador se dirige, vou passar para o nosso próximo debatedor, apresentador, que é o – deixe-me olhar aqui; não consigo enxergar mais a essa distância – Rodrigo Badaró.

O Dr. Rodrigo Badaró é Conselheiro Nacional do Ministério Público.

Doutor, o senhor tem dez minutos para a sua apresentação.

Obrigado.

O SR. RODRIGO BADARÓ (Para expor.) – Senador, é uma alegria estar aqui, nesta Casa, nesta nessa Câmara Alta. Deixo meu abraço aqui ao Senador do meu estado, Senador Izalci, um abraço ao Senador Eduardo Gomes, que nos deixou.

Senador, antes de mais nada, além de ser Conselheiro Nacional do Ministério Público, sou advogado e represento a OAB Nacional também, como Presidente da Comissão Nacional de Proteção de Dados, que antes era presidida pela minha querida amiga Estela, que tão bem exerce o papel hoje no Ministério da Justiça. A OAB emprestou a Dra. Estela para desenvolver um belíssimo trabalho no Ministério da Justiça.

(Intervenção fora do microfone.)

O SR. RODRIGO BADARÓ – Exatamente.

E aqui meu papel, ao falar depois de tantos especialistas e de um debate tão complexo, Senador, é, primeiro, parabenizar a capacidade regulatória e legislativa e a Comissão de Juristas, presidida pelo



SENADO FEDERAL
Secretaria-Geral da Mesa

Ministro Cueva, que apresentou um trabalho técnico relevante, e parabenizar mais ainda, Senador, o trabalho de V. Exa. e desta Comissão de desenvolver o trabalho e dar ao povo, aos técnicos abertura para discutir. Por quê? Em 30, 40 minutos aqui, nós notamos que, não obstante estarmos falando de máquina e de tecnologia, graças a Deus, a preocupação é com o ser humano. E, ao mesmo tempo, há uma preocupação, colocada pela Profa. Fernanda, com a questão racial, que tem que ter, evidentemente, e uma preocupação, muito bem colocada pelo Prof. Parentoni, com a liberdade econômica, com a inovação, com o empreendimento da inovação. E todos esses são princípios constitucionais. Então, quando a gente fala de princípio, a gente está falando do princípio da dignidade humana; o Parentoni está falando do princípio da liberdade econômica, da liberdade de expressão, da livre atividade econômica, do desenvolvimento da inovação.

Eu penso que, ao analisar essa legislação, Senador, e aqui sem qualquer visão de regulamentação ou não – e concordo com V. Exa. quando diz que há naturalmente um estado de neofobia, o medo, realmente, da inovação –, me preocupa o ímpeto regulatório e as interpretações e limitações que isso pode gerar.

O que eu acho que é importante? Quando a gente fala aqui de princípios éticos, e a Profa. Fernanda colocou tão bem com relação à diferenciação do Brasil, a ética nada mais é que um trecho filosófico de problemática cultural de uma sociedade. Os nossos valores éticos são certamente diferentes dos valores dos europeus.

E aqui eu brinco, Senador, como fiz outro dia, quando estava dando uma palestra: eu falo que o europeu – e nós todos, querendo ou não, temos um quê de europeu – não tem talvez tanta capacidade tecnológica, tanta capacidade militar, mas tem uma capacidade de imposição moral absurda. Então, a gente tem que ter um cuidado, porque o brasileiro – e, nesse ponto, não podemos ter esse espírito e nos apequenarmos – tem um país continental e uma questão ética e problemas que nenhum lugar do mundo tem. Então, é importante, quando da interpretação da norma, quando do desenvolvimento da norma, termos em mente essa questão específica brasileira, como tão bem colocou a Profa. Fernanda quanto à questão racial; e como tão bem colocou o Parentoni: será que nós temos hoje capacidade de impor a pequenas empresas, já com carga tributária elevada, Senador, regras e regras que impeçam o desenvolvimento em tecnologia? Todos sabem que o dinheiro do desenvolvimento da tecnologia é complexo e é volumoso.

E, dentro desse espectro, aqui a minha maior preocupação, do ponto de vista de advogado... E eu me sinto um ser híbrido no sentido institucional, eu brinco com isso – e muito orgulhoso de ter sido até indicado por esta Casa e de ter tido meu nome referendado por esta Casa para o Conselho Nacional do Ministério Público –, porque o Ministério Público hoje e a OAB são reflexos, são instituições que refletem muito o sentimento do cidadão, porque elas representam, de certa forma, no espectro judicial, a cidadania. Então, isso é importante. E qual é a preocupação sempre do Ministério Público e da OAB? Segurança jurídica. Então, a grande pergunta que eu coloco aqui no debate regulatório é: que segurança jurídica, sem um cerceamento de amplitude de direitos e sem afrontar princípios, nós daremos para o cidadão? É possível colocar isso num texto?

Nossa Constituição já provou que não é, Senador. Nós temos uma Constituição tão literal e tão extensa... mas olha o tanto de mazelas que nós temos no Brasil. Será que regular tanto, que colocar tanto em um texto vai trazer segurança jurídica? Será que uma autorregulação e uma autocontenção setorial... E a Profa. Estela sabe bem que o direito e a lei, no nosso sistema de *civil law*, só são gerados após um dano. Então, será que não é melhor ampliar o debate e esperar os danos? Por mais que isso seja, do ponto de vista de atonicidade, do ponto de vista dramático... Porque o direito é criado após



SENADO FEDERAL
Secretaria-Geral da Mesa

uma litigância de interesses e uma contenção e um problema de dano. Lembro aqui que a pretensão resistida de um ente público ou privado é que gera um processo.

Então, eu queria trazer aqui, de forma muito clara, a preocupação no sentido de tentar... E aqui eu não estou, Presidente...

Parabenizo aqui também a minha querida amiga Patrícia Peck, que trouxe argumento, a transparência, os princípios; e a Dra. Estela nos brindou com outras questões de direito subjetivo que hoje estão em debate, até do ponto de vista global.

Os princípios estão já na nossa Constituição. Eles estão lá. Não é melhor a gente tentar interpretá-los de uma melhor forma dentro da tecnologia? Porque, como bem colocou a Profa. Peck – e todos nós – e o Senador acabou de citar, a tecnologia não negocia com tempo, ela não negocia com o Parlamento, ela não negocia com norma. Então, há de se desenvolver um pensamento.

E aqui, já antecipando, até sou a favor de uma conciliação regulatória, Senador. Eu acho que hoje a gente tem um avanço regulatório enorme na Anatel. A Anatel é um órgão hoje robusto. Nós temos já uma nova ANPD. E o Senador Astronauta Marcos Pontes colocou bem: não tem como ter inteligência artificial sem ter dados, então não tem como desvincular a Autoridade Nacional de Proteção de Dados desse debate. E aqui eu não quero antecipar que eu tenho preferência de que a regulação seja da Anatel ou seja da ANPD ou que se crie uma nova autoridade, que é o que prevê a lei – talvez a possibilidade de criação de uma nova autoridade. E lembro, Senador Izalci, que hoje o Governo discute também a criação, salvo engano, de uma agência de cibersegurança vinculada mais à questão de segurança nacional.

Então, nem querendo ser como Chomsky, que não acredita muito que sequer há inteligência artificial, nem também como um pensador que eu gosto de citar muito, que é o Thomas Kuhn, que escreveu em 1962 a questão da estrutura tecnológica, da quebra de paradigma, eu não sei se nós estamos dentro dessa evolução tecnológica extraordinária que o Prof. Thomas Kuhn falava. Será que estamos? Será que essa inteligência artificial é um risco fabricado ou é um risco natural? Há de se questionar isso também. Quais os interesses econômicos por trás da inteligência artificial? Qual o medo? Quais as consequências e as limitações que, ao impor às empresas, à sociedade, ao Judiciário, serão auferidas desse questionamento?

(Soa a campainha.)

O SR. RODRIGO BADARÓ – Então, aqui, Senador, estou provocando mais reflexões do que ajudando em qualquer esclarecimento.

Tinha um professor meu, que eu até citei ontem na sessão do CNMP, que falava que, quando o advogado não consegue esclarecer muito, ele pelo menos tem que confundir. Então, acho que a confusão gera o debate, e a mistura de pensamentos gera talvez a solução, que vai ser certamente a melhor para o Brasil. E não tenho dúvida de que esse ambiente de debate aqui promovido pelo nosso Senado Federal é o campo ideal para o ambiente qualificado, dando voz à sociedade e a todas as pessoas de conhecimento técnico que aqui hoje estão presentes.

Muito obrigado, Senador.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Obrigado, Dr. Badaró. Sem dúvida nenhuma sua contribuição é muito importante.

Um ponto que eu gostaria de citar é que, dentro dessas reuniões que nós temos tido aqui, algumas coisas estão de certa forma se cristalizando. E, logicamente, eu vou levar tudo isso ao Relator, o Senador Eduardo Gomes, e também ao Presidente da Comissão.



SENADO FEDERAL
Secretaria-Geral da Mesa

Por exemplo, para tipos de tecnologias ou tipos de áreas que são completamente transversais, como inteligência artificial, como segurança cibernética, para citar essas duas, fica muito difícil estabelecer – isso na minha opinião e na opinião de muitos dos debatedores que aqui vieram – uma agência única capaz de interferir. Isso vai causar interferências com outros, porque a aplicação é setorial. Não é como se você falasse assim: "Olhe, eu vou definir aqui a parte de espaço". Vamos chamar assim: espaço. Eu crio uma agência espacial que cuida de espaço, e fica relativamente claro e simples que o foco está aqui dentro.

Uma agência tem as suas funções, como agência reguladora, de justamente trabalhar com a regulação, a fiscalização, a aplicação de multas, a aplicação daquela lei, etc. Como isso é uma aplicação setorial, fica simples você pensar nessas legendas, mas, quando você fala em segurança cibernética, que está em todos os setores, está no setor crítico da energia, por exemplo, está em setores mais simples de aplicações no dia a dia, está em todos os lados, fica difícil você estabelecer, na minha opinião e na opinião de muitos aqui, uma agência que cuide disso. Mas é possível se fazer – é uma das soluções pensadas, e a gente vai propor esta solução – um conselho composto por membros de cada um desses setores, das agências que já têm essa funcionalidade, essa ação de regular e fiscalizar cada um dos setores, e cada um deles vai ter ali o seu setor de segurança cibernética, o seu setor de inteligência artificial. O composto desses conselhos pode estar centralizado, por exemplo, no GSI ou em algum outro órgão já existente. Aí, sim, toma-se a decisão baseada na visão ou na perspectiva de cada um desses setores, porque uma agência ter pessoas que conheçam todos os setores, para regular todos os setores, além de ter a interferência dentro dos setores, fica muito difícil de funcionar na prática. E a gente tem que pensar em coisa prática. A possibilidade de se ter isso tem se transformado aqui bastante, tem se cristalizado bastante isso aí. Então, vamos ver como é que isso se reflete nos resultados. Logicamente vai depender do Relator e de tudo mais, mas é um dos pontos que nós estamos levando para ele.

Eu gostaria de aproveitar também e passar a palavra ao Senador Izalci para suas considerações.
Por favor, Senador.

O SR. IZALCI LUCAS (Bloco Parlamentar Democracia/PSDB - DF. Para interpelar.) – Bom dia ainda, não é?

Primeiro, quero cumprimentar o Astronauta Marcos, nosso Senador, ex-Ministro dessa área e bastante conhecedor da área de ciência, tecnologia e inovação, por essa iniciativa.

Quando se começou a discutir a questão da inteligência artificial, a gente discutiu muito aqui o medo dessa regulamentação muito rápida. O Brasil gosta de regulamentar tudo. A gente nem sabe o que é ainda, como é que vai ser isso, e o Brasil já está regulamentando e impedindo até ou dificultando o avanço da inovação.

E a gente tem que encarar o mundo real. O mundo real hoje é que as coisas vão avançar, independentemente de a gente querer ou não, porque hoje não é mais... A questão não é local ou regional, a questão é global. E você vê que os países desenvolvidos nessa área ainda têm muito pouca regulamentação, exatamente porque há um espaço muito grande.

A minha maior preocupação – e a gente precisava organizar isso – é a questão da educação. Nada se faz sem educação. A gente vê cada vez mais a nossa educação totalmente desconectada do mundo real da inovação. E eu estive agora naquele Congresso de Inovação em São Paulo e fiquei muito preocupado, porque quanto mais inovação, quanto mais conhecimento, mais vão aumentando as distâncias entre aqueles que mais podem e aqueles que menos podem. E houve um esforço danado aqui para a gente fazer o novo ensino médio para os jovens começarem a ter uma profissão ou entrar nessa área de tecnologia, mas... Foram cinco anos para a gente começar um projeto, dois anos já



SENADO FEDERAL
Secretaria-Geral da Mesa

iniciados, e agora vem uma proposta totalmente... São três passos atrás, nem vou dizer dois. Então, isso me preocupa muito, porque não tem outra solução que não seja compatibilizar tudo isso com a educação de base, não é a educação superior. A gente tem que começar na educação infantil e no ensino fundamental. E hoje eu fico triste realmente, porque a gente não vê nessas ações um caminho de melhorar a educação.

E aí me preocupou muito e também o Senado Marcos, quando o Senado criou esta Comissão, com aspecto muito mais jurídico. A gente tem aqui... Eu até parabeneizei os advogados... Eu vi agora a reforma tributária dizendo que só pode participar do Consei se for advogado. Os advogados têm um poder muito forte de convencimento aqui na ordem nacional. E aí já queriam começar, como começaram, a tratar da inteligência artificial já num foco jurídico.

Olhem: como fazer essa regulamentação? A gente não sabe nem o que é isso ainda, agora que está se desenvolvendo. Então, por isso, nós propomos várias audiências. E aí foi criada, inclusive, Senador Marcos Pontes, esta Subcomissão na Comissão de Ciência e Tecnologia para a gente aprofundar um pouco mais. Pelo menos o meu objetivo principal é ouvir um pouco mais, porque é uma coisa tão ainda embrionária, vamos dizer assim, que vai avançar rapidamente, que a gente precisa ouvir todo mundo e tentar compatibilizar tudo isso. A gente tem dificuldade nas *startups*, para as novas empresas, os jovens não têm apoio. Até pouco tempo, 80% da energia consumida eram na burocracia, nessas partes que não têm nada a ver com o foco. As pessoas, às vezes, têm uma ideia maravilhosa, uma condição de desenvolver muito, mas aí ficam presas nessas questões burocráticas. Então, a gente está ouvindo aqui.

Evidentemente, nós vamos debater este tema ainda por algum tempo para a gente poder botar em pauta e votar. Há uma preocupação muito grande, principalmente no segmento de tecnologia, e, ao mesmo tempo, a gente fica preocupado. Eu vi agora, ontem, antes de ontem, já uma discussão... ações judiciais nos Estados Unidos com relação a neto e outras sobre a questão da criança, do vício. Porque se deixar, realmente, as netinhas, hoje, os netinhos, aí são... Já querem... E vô não pode nem falar não; vô faz... E a mãe às vezes quer se livrar do negócio e já passa o celular para ela; fica 24 horas por conta disso. Então, há umas coisas com que a gente fica preocupado, mas, da mesma forma, a gente fica preocupado também de não ficarmos para trás numa evolução muito grande que está acontecendo no mundo todo.

E aí essa regulamentação equivocada pode trazer consequências gravíssimas em termos de tecnologia de inovação. Então essa é a preocupação, Marcos. A gente está aqui, luta, e essa área vocês sabem que não é fácil. Ao se falar em ciência, tecnologia e educação no Brasil, a gente tem apoio total, mas só no discurso, não é? Na prática mesmo, quando se vai botar recurso, etc., aí fica para depois, não é?

É isso. Quero parabenizar mais uma vez. Estamos aqui ouvindo. Vamos olhar depois as apresentações todas. Eu estava na reunião de Líderes, perdi algumas aqui, mas é um tema que nós vamos acompanhar de perto, para não regulamentar demais e nem deixar solto, não é?

Parabéns, Senador Marcos.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Obrigado. Obrigado, Senador Izalci. Sempre participante. É uma das pessoas aqui no Congresso que sempre está ao lado do desenvolvimento da ciência no país.

Enquanto ele falava, uma coisa que me passou... uma comparação. É um pouco simples, mas às vezes é importante para as pessoas perceberem, em termos dessa legislação de regulação de setores.



SENADO FEDERAL
Secretaria-Geral da Mesa

Acho que todo mundo já teve a experiência de ter de colocar um antivírus no seu computador. Você coloca um antivírus no seu computador, obviamente com a expectativa de bloquear algum *malware*, algum tipo de *software* ruim para entrar no seu computador. Algumas vezes você instala um antivírus que é muito bom, mas muito bom mesmo, mas aí para o seu computador, você não consegue mais fazer absolutamente nada no computador; fica travado, lento e você não consegue... Agora, às vezes, você coloca um que – você está fazendo tudo – não bloqueia muita coisa.

O ideal é ter um antivírus que permita que o seu computador opere na velocidade normal, com todas as necessidades que você tem, e, ao mesmo tempo, bloqueie os ataques principais, cibernéticos, aí na sua máquina.

Eu vejo essa legislação mais ou menos como isto também: uma maneira de bloquear as coisas ruins, mas, ao mesmo tempo, deixar que tudo seja feito. E ninguém consegue prever o que você vai fazer no seu computador, você pode estar desenvolvendo até uma inteligência artificial ali dentro ou simplesmente usando o texto. Mas esse antivírus tem que prever que há coisas inesperadas que você vai fazer com aquele computador.

Agora, na sequência das nossas apresentações, eu gostaria de passar a palavra para a Dra. Cynthia Picolo, Diretora-Presidente do Laboratório de Políticas Públicas e Internet (Lapin).

Dra. Cynthia, você tem os dez minutos. E eu peço para controlar o tempo por lá, porque não vai ter as indicações por aqui.

Obrigado.

A SRA. CYNTHIA PICOLO (Para expor. *Por videoconferência.*) – Primeiramente, bom dia. Eu gostaria de começar agradecendo esse convite da Comissão, nas pessoas dos Srs. Senadores Carlos Viana, Marcos Pontes e Eduardo Gomes, e também estendo os meus agradecimentos à Adriana, pela abertura de sempre para diálogo com a sociedade civil.

Eu sou a Cynthia Picolo, Diretora do Lapin, que é uma organização de pesquisa e ação sediada em Brasília, e somos integrantes também da Coalizão Direitos na Rede. A minha exposição aqui, neste breve tempo, toca em quatro aspectos principais, que são os aspectos ambientais, laborais, de preservação cultural e de governança multissetorial.

Então, para começar abordando questões ambientais, é preciso que a gente tenha consciência de como a tecnologia vem impactando o meio ambiente. Além da fala da pessoa indígena de ontem, que foi o Time'i Assurini, eu adiciono aqui alguns dados baseados em estudos, como, por exemplo, o de que a emissão de dióxido de carbono, durante o ciclo de treinamento de grandes modelos de IA, corresponde a quase cinco vezes as emissões de vida útil de um carro médio, ou, então, que um *data center* em construção no Uruguai consumirá até 6, 7,6 milhões de litros de água por dia, nos dias de verão; isso corresponde a um consumo diário de água potável de 55 mil pessoas. Além disso, eu relembro aqui o caso envolvendo o garimpo e a compra ilegal de ouro ianomâmi por uma fornecedora de *big tech*.

Então, a partir desse contexto e desses dados concretos, o que pensar para o projeto de lei para o novo instrumento regulador de IA no aspecto ambiental? Precisamos de transparência e prestação de contas. O PL deve trazer obrigações específicas para o fornecimento de informações acessíveis e compreensíveis a respeito do impacto ambiental para treinamento e funcionamento dos sistemas e deve ter, no mínimo, informações sobre consumo de energia, sobre equipamentos usados e onde são hospedados os dados, para, então, permitir o mapeamento de consumo de recursos minerais e água. É preciso também mapear *data centers*, para isso auxiliar o processo de auditoria para compreensão sobre como estão funcionando e participando do treinamento dos sistemas.



SENADO FEDERAL
Secretaria-Geral da Mesa

Para além disso, deve haver uma prerrogativa do regulador para que ele consiga determinar, com base em consulta à comunidade técnica e sociedade civil, parâmetros e limites para o consumo de recursos naturais, para o desenvolvimento e uso desses sistemas, porque somente com a definição desses parâmetros é que o regulador poderá ter poder de agência nos aspectos ambientais, e ele sai, portanto, da esfera de apenas fornecimento de documentação. Então, vale lembrar, nesse ponto, que a avaliação de impacto algorítmico, que felizmente já está prevista no PL 2.338, é um instrumento importante no mapeamento do impacto ambiental e de salvaguardas e deve incluir as informações que eu acabei de citar sobre transparência a respeito de consumo de recursos. E eu relembro aqui que tudo está alinhado aos objetivos de desenvolvimento sustentável da ONU, e é importante fazer também essa correspondência.

O segundo ponto que eu quero tocar é sobre as questões laborais. Há diversos dados, hoje, sobre o aumento da plataformização do trabalho, que é um cenário permeado por desigualdades e, no contexto da IA generativa, por exemplo, a gente já sabe que grandes empresas usam mão de obra para realizar tarefas de catalogação de dados e moderação de conteúdo mediante remuneração e condições de trabalhos degradantes. Então, além de ser financeiramente insatisfatório, esse tipo de trabalho afeta a saúde física e mental das pessoas, já que elas são submetidas, por exemplo, a analisar conteúdos ofensivos e violentos.

Por outro lado, há trabalhadores que dependem de sistemas algorítmicos para exercer as suas funções, como no caso a plataforma Uber, em que o sistema classifica motoristas e determina as corridas de acordo com a disponibilidade, a localidade e outros fatores.

Ontem mesmo saiu uma pesquisa da Pnad que revelou que, em 2022, o Brasil tinha 1,5 milhão de pessoas que trabalhava por meio de plataformas digitais e aplicativos de serviços. Assim, sabendo dessa discriminação e opacidade que permeiam os sistemas automatizados, toda forma de gerenciamento de trabalho via algoritmo deve ter participação dos trabalhadores no pensar desse sistema, principalmente na questão de determinação de parâmetros.

Portanto, eu volto a reforçar a necessidade de transparência, a transparência sobre o trabalhador ter pleno conhecimento dos parâmetros do sistema para fazer esses *outputs*. E, além disso, é preciso pensar na ingerência dessas pessoas, eventualmente por meio de um conselho, para que decisões tomadas pelo e sobre o sistema tenham participação desses trabalhadores.

Vou encaminhando para o final para endereçar esses dois últimos pontos, e eu gostaria de falar sobre diversidade cultural e governança multissetorial.

Pensando um pouco na proteção da diversidade cultural, nós vimos ontem também na fala do indígena, do Time'i, a importância de se preservarem as várias formas de expressões culturais. Então, pensando na comunidade indígena, por exemplo, e também trazendo dados, segundo o IBGE, em 2021, o Brasil contava com mais de 800 mil pessoas indígenas, com cerca de 274 idiomas nativos. Então, a IA tem que preservar essas linguagens. Assim, a gente precisa propor mecanismos e pensar neles para que essas expressões culturais, bens artísticos, históricos, materiais ou imateriais, sejam protegidos e incorporados no contexto da IA. Caso eles não sejam, há um risco gigantesco de apagamento cultural que pode afetar diretamente a autodeterminação dessas pessoas. E eu relembro, lógico, que a comunidade indígena é apenas um exemplo, temos várias outras diversidades culturais aqui no nosso país.

E, por fim, eu trago uma pequena reflexão também sobre a questão de governança multissetorial no ecossistema da IA. A gente está ainda nesse quebra-cabeça sobre entidade reguladora, que é algo natural nessa fase de reflexão sobre o que a gente quer em termos de IA no Brasil e monitoramento,



SENADO FEDERAL
Secretaria-Geral da Mesa

coordenação, mas, mesmo que a gente não tenha essas respostas prontas, é importante reforçar a necessidade de coordenação e articulação entre os reguladores setoriais, porque, certamente, cada regulador conhece melhor o seu setor, e as suas prerrogativas nesse contexto não devem ser diminuídas. No entanto, é preciso pensar nesse arranjo de coordenação entre esses reguladores setoriais para criar segurança jurídica para as empresas e para as pessoas. E, nesse contexto, não haveria uma invasão de competência dos reguladores setoriais, mas, sim, uma orquestração para lidar com a IA nesses diversos setores.

Então, é preciso, sim, que haja uma entidade central que faça esse papel de coordenação e articulação conjunta, porque a gente está falando de IA, que é uma tecnologia disruptiva, que vai continuar mudando e afetando a sociedade, o mercado, e que vai exigir novos estudos e novas abordagens. Então, sem uma entidade coordenadora para definir estratégias, prioridades e estudos, em conjunto sempre com esses setores, a gente vai ter um cenário fragmentado. E esse cenário fragmentado já foi, inclusive, diagnosticado na investigação que está em curso para a aplicação da Metodologia de Avaliação de Prontidão da Unesco, a qual o Lapin está apoiando.

Então, pensando nisso, os órgãos reguladores poderiam, por exemplo, criar grupos de trabalho com os mais diversos especialistas, e esses grupos teriam, então, a função de conversar e oferecer subsídios específicos daquele grupo de fora à entidade coordenadora.

E, chegando ao final da minha exposição, eu reajo sobre um ponto: o de a lei não regular oportunidade. De fato, não é objetivo da lei, até porque oportunidade é questão de política pública, e a gente tem a Ebia, que pretende melhorar e tem as ações estratégicas para fomentar o setor produtivo. E o PL 2.338, inclusive, traz o *sandbox* regulatório como um instrumento para fomentar a inovação no país. Então, diante desse contexto, eu ressalto que a abordagem não pode ser somente principiológica. A nova lei deve dar contornos à efetivação dos princípios, direitos e responsabilização.

E eu fecho a minha fala também fazendo coro às mais diversas vozes sobre a necessidade de debate inclusivo, participativo, de capacitação digital, incluindo incentivo a meninas e mulheres... Então, tem muitas coisas aí para a gente debater e pensar em conjunto.

Também me coloco à disposição para me somar no debate. E o Lapin está sempre presente, apoiando os gestores públicos. É para a gente uma grande honra utilizar esse espaço para essa conversa. Muito obrigada.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Obrigado, Dra. Cynthia Picolo, Diretora-Presidente do Laboratório de Políticas Públicas e Internet (Lapin).

A próxima debatedora, apresentadora é a Dra. Ana Carla Bliacheriene, Professora da Universidade de São Paulo.

Professora, a senhora tem dez minutos.

A SRA. ANA CARLA BLIACHERIENE (Para expor.) – Bom dia a todos. É um prazer estar aqui.

Quero agradecer o convite que me foi feito para estar nesta Comissão pelo Senador Eduardo Gomes e à Adriana, que o apoia e sempre nos acolhe.

E quero dizer da minha satisfação e alegria de estar numa sessão presidida pelo meu Senador Astronauta Marcos Pontes, Senador do meu estado de adoção – eu adotei o Estado de São Paulo há 25 anos, e o Estado de São Paulo me adotou. E o meu estado de nascimento, Sergipe, eu muito amo também, mas, em São Paulo, eu desenvolvi família, me desenvolvi profissionalmente e lá exerço as minhas atividades na Universidade de São Paulo. E, como Professora universitária, advogada também, mas também como cientista, quero dizer do meu prazer de tê-lo nesta Comissão, porque acompanho



SENADO FEDERAL
Secretaria-Geral da Mesa

já há algum tempo a luta de V. Exa. no sentido de valorizar a ciência, a tecnologia e o que se produz dentro das universidades. Sempre que pôde, foi visível o esforço que V. Exa. fez, desde o início do mandato, para que a ciência fosse respeitada neste país e para que pudesse se desenvolver. Então, é uma honra muito grande estar numa mesa presidida por V. Exa.

O tema da nossa audiência é um tema que abrange muitos assuntos, muitos temas. E eu resolvi abordar três com maior enfoque, que são: a questão da regulação baseada em risco; as decisões automatizadas; e a supervisão humana. Então, dividida em seis partes: uma pequena apresentação do que fazemos na universidade; vou falar um pouco sobre o PL; a regulação baseada em risco; os benefícios para o setor público do uso da IA; e as decisões automatizadas, se houver tempo.

Então, sou Professora da Universidade de São Paulo, advogada de formação; trabalho com os temas do direito e tecnologia; coordeno o grupo de pesquisas *smart cities* BR na Universidade de São Paulo. Tive o prazer de fazer parte da primeira versão da comissão do CNPD, junto com o Leonardo, com a Peck, a quem eu cumprimento e parabenizo pelas falas; e também coordenei o Comitê de Inovação e Transição Digital de Governos do Instituto Rui Barbosa, braço acadêmico dos tribunais de contas do Brasil.

Falo aqui em nome de um grupo de pesquisadores, coordenado por quatro professores da Universidade de São Paulo: eu; o Prof. Luciano Vieira de Araújo, que teve oportunidade de falar ontem; o Prof. Dib Karam Junior; a Profa. Fátima Nunes. É um grupo interdisciplinar onde trabalham professores da área do direito, gestão pública, engenharia, ciência de dados e da computação, além de vários outros profissionais, porque o tema das cidades inteligentes é um tema interdisciplinar.

Para fazer a preparação para esta audiência pública, nós não só falamos de IA como usamos massivamente IA no nosso grupo de estudos na universidade. Então nós processamos todos os documentos produzidos, todos os vídeos e, a partir desses documentos que processamos e desses vídeos, geramos resumos de documentos, listamos os principais pontos abordados, comparamos documentos e listamos diferença entre documentos.

Acabei de falar para o Senador Eduardo Gomes que esse sistema está à disposição desta Comissão para gerar relatórios que sejam necessários para a resolução, o relatório final do trabalho.

Um pouco sobre o PL. Primeiro eu quero fazer aqui uma referência à importância do trabalho não só dos juristas do PL 2.338, de 2023, mas também de todos os que trabalharam desde 2019 oferecendo visões e opções, além dos servidores desta Casa que têm feito um trabalho belíssimo.

A função destas audiências públicas, de fato, é avaliar aquilo que é positivo ou não. Num breve passar do projeto de lei, eu vou aqui ao art. 1º, que fala que o objetivo é proteção dos direitos fundamentais e garantir a implementação de sistemas seguros e confiáveis. Leis não garantem a implementação de sistemas seguros e confiáveis; políticas garantem; e práticas garantem. Então a lei não tem capacidade de entregar o que está dizendo.

Eu ficaria feliz se eu visse nos objetivos de um projeto final que o objetivo fosse a proteção dos direitos fundamentais e o desenvolvimento científico e tecnológico do país. Penso que, numa lei que trate de IA, o objetivo deve ser esse. E, dentro desse grande guarda-chuva, a gente vai ter o benefício da pessoa humana, o benefício da democracia e a proteção dos direitos fundamentais, que eu senti no projeto de lei sobrevalorizada como se fosse o capítulo não escrito da LGPD. Eu senti que o projeto era como se fosse o capítulo que faltou na LGPD. E, embora sejam temas muito aproximados, a proteção de dados é só uma parte da proteção do indivíduo, ela não cobre a proteção do indivíduo, e o tema da IA é muito maior do que o tema da proteção de dados.



SENADO FEDERAL
Secretaria-Geral da Mesa

Então, eu percebi confusão e redundância entre fundamentos e princípios; comprometimento da inovação, desenvolvimento econômico e competitividade nacional; sobrevalorização da supervisão humana, mesmo em situações em que a IA atuaria de forma melhor e menos custosa – o Leonardo já trouxe as situações e as classificações –; não é capaz de regular todas as modalidades da IA – também falado pela PEC e pelo Leonardo –; confusão com conceitos da LGPD e do Código de Defesa do Consumidor, fazendo paralelos não adequados; algumas partes parecem o capítulo não escrito da LGPD; ficção de que tudo em IA pode ser traduzido num caminho, fórmula ou algoritmo 100% rastreável – isso é típico de uma comissão que não tinha representantes da área científica própria, não é?; é um erro terminológico do conceito da IA –; aplicação de IA em setores estratégicos, como indústria, saúde, segurança, educação e agricultura, fica comprometida.

Há uma diferença entre regulação baseada em risco e regulação baseada no medo e na punição. A sensação que eu tive do projeto de lei foi a de um projeto regulado no medo e na punição e não no risco. E por que eu digo isso? Qualquer regulação baseada em risco passa por um método, por uma técnica. Aqui não estão todas as etapas de análise de risco – coloquei apenas poucas –: identificação clara do risco associado ao tema, coleta de dados relevantes, análise quantitativa e qualitativa dos riscos identificados, escolha do modelo da avaliação, comunicação do risco. Depois disso, a gente passa para estratégias de mitigação e, aí, sim, para proposta legislativa.

O que é que eu percebi? Embora se use a terminologia "regulação baseada em risco", que é importantíssima e, se o Senado acolher isso, será relevante, não foi utilizada essa técnica para a realização do projeto de lei ou pelo menos a técnica que recomendam os modelos internacionais de risco. E, aí, o que faz? Eu falo aqui em alguém que estuda aplicação de IA no setor público. Quais são os benefícios? Há uma pergunta aqui: quais os benefícios de IA no setor público? São 100% de benefício. Os três Poderes precisam não é de uso de IA, mas de uso massivo de IA para qualificar o serviço público. Estamos discutindo uma reforma tributária infundável há décadas e a contraparte dela é a reforma administrativa. Talvez a grande revolução da reforma administrativa ainda não feita seja a utilização massiva de IA na administração pública, que vai lhe conferir maior eficiência ao menor custo. E aí a gente teria respostas para a reforma tributária.

Então, áreas como infraestrutura crítica, recrutamento e gestão de trabalhadores, serviços públicos essenciais e administração da justiça necessitam fortemente da atuação da IA ou nos transformaremos em neoludistas, ludistas da quarta revolução industrial. Já tivemos o da primeira, seríamos neoludistas. O tempo será inclemente para as nações e para as organizações públicas e privadas que adotarem esse caminho.

Então, entendo que regular, sim, e avaliar riscos muito claramente quanto a direitos fundamentais e preservação da centralidade humana, mas não regular à base do medo, que foi o que eu senti avaliando a legislação, embora ela tenha pontos relevantíssimos e bastante modernos.

Para o Poder Legislativo, a IA cria um processo legislativo baseado em evidência; análise de legislação com comparação legislativa, inclusive nos três níveis federais, há um problema grave no Brasil de sobreposição legislativa; engajamento cidadão, com participação ativa no processo legislativo; automatização de tarefas administrativas do processo administrativo; transparência ampla do processo administrativo; e monitoramento, que existe – inclusive eu quero fazer aqui uma referência a uma servidora desta Casa, Alba Valéria, que fez um trabalho excelente sobre avaliação de políticas públicas a partir do Poder Legislativo –, aqui a IA daria o monitoramento e indicadores precisos sobre o impacto prévio e posterior de leis e de políticas públicas. Isso um modelo humano não será capaz de nos dar com qualidade.



SENADO FEDERAL
Secretaria-Geral da Mesa

No Poder Executivo, aprimoramento do serviço público com melhora de eficiência, eficácia, efetividade, economicidade e, pasmem, equidade, garantindo não só não discriminação, como inclusão de pessoas necessitadas através da IA; tomada de decisão baseada em evidência ou política pública baseada em evidência é o grande novo tema que se estuda na academia para a transformação do setor público; automatização de tarefa; personalização de serviços governamentais, atendendo a necessidades muito próprias do indivíduo; gestão de recursos humanos para a tal da grande reforma administrativa.

No Poder Judiciário, é essencial que se retire do projeto a colocação de atuação no Poder Judiciário como risco com todas aquelas implicações para o ecossistema de inovação e de *startups* no Poder Judiciário. O Poder Judiciário está asoberbado. Existem ações massivas entrando diariamente. Há pessoas que entram com 40 ações no mesmo dia, com o mesmo objeto partilhado, para ver onde que vai cair melhor.

(Soa a campanha.)

A SRA. ANA CARLA BLIACHERIENE – São as ações predatórias e ações massivas sem o uso de IA, que os advogados... Nós advogados já usamos IA nos nossos escritórios. Como que o juiz não vai ter a IA para apoiá-lo em processo de classificação, processos de leitura de relatórios, processos de sugestão de minutas? Isso facilita muito. Em alguns casos específicos, eu defendo, inclusive, o juiz robô, que é algo que o CNJ por enquanto não regulamentou, mas são casos muito específicos.

Automatização de processo; análise jurisprudencial; geração de documentos; detecção de fraudes e corrupção; acesso à Justiça facilitado e diminuição do tempo do julgamento, que é uma chaga do Poder Judiciário; monitoramento de execução de sentenças e expedição de mandados. Um dia em que um réu fica a mais no sistema prisional, e não deveria, é um dia de injustiça. A IA resolve isso.

(Soa a campanha.)

A SRA. ANA CARLA BLIACHERIENE – Se alguns colocam que a IA não deve ser usada na área criminal porque ela encarcera, eu digo que a IA também desencarcera no dia e na hora certa, sem que um cidadão tenha de ficar um dia a mais no sistema prisional.

Decisões automatizadas, um tema importante, tem bastante restrição dentro da lei. Eu não sou "Poliana", eu compreendo o porquê de todas as restrições, não há espaço nesse tempo que nós temos para discutir, mas, Senador, eu me refiro à análise de risco e benefício. Eu não estou falando custo-benefício, eu falei risco-benefício. Quando falamos de direito fundamental, não é custo, não há custo, é risco-benefício. E, na análise de risco-benefício dentro do setor público, dos três Poderes, não há dúvida de que o uso massivo da IA será uma transformação importante na qualificação da entrega dos três Poderes.

(Soa a campanha.)

A SRA. ANA CARLA BLIACHERIENE – Por fim, eu digo que, nesse modelo de conselho e agência que V. Exas. estão pensando, considerem não só os setores, porque aí a gente estaria regulando a área do mercado, mas a representação do chefe dos três Poderes, porque, no setor público, o uso massivo da IA é a solução para a gente melhorar a entrega do serviço público e, quem sabe, um dia ter um país com menos desigualdade, porque pagamos o imposto, recolhemos muito, somos a décima economia do mundo – a décima economia do mundo não é pobre, nos engana quem diz que somos pobres. A gente precisa melhorar o serviço público para que todos tenham uma vida decente e digna.



SENADO FEDERAL
Secretaria-Geral da Mesa

Muito obrigada.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Excelente. Muito obrigado, Profa. Ana Carla, parabéns pela apresentação.

Na sequência, eu gostaria de passar a palavra para... Vou achar aqui. Pode ser para a Dra. Adriana Rollo, Líder da Comissão Especial de Regulação de IA da Associação Internacional de Inteligência Artificial.

A senhora tem dez minutos, por favor.

A SRA. ADRIANA ROLLO (Para expor.) – Obrigada, Senador. Em nome da Associação Internacional de Inteligência Artificial, agradeço muito o convite do Senador Eduardo Gomes para estar aqui, na presença do Senador Astronauta Marcos Pontes.

A associação é muito comprometida com a agenda positiva da IA, por isso eu faço coro aos debatedores que vieram antes de mim, porque me parece que nós temos um uníssono aqui em relação à importância de desenvolvermos a indústria tecnológica de IA em benefício do Brasil.

Eu queria começar a minha fala lembrando que estamos em outubro, e aqui nós temos o Outubro Rosa, uma campanha muito importante de ações afirmativas de incentivo e prevenção de câncer de mama, que é uma doença que mata mais de 230 mil mulheres por ano, é o câncer mais letal entre mulheres, e por isso eu queria aproveitar o exemplo prático de uma *startup* brasileira que desenvolveu um sistema de IA para fazer a leitura dinâmica de laudos e exames médicos, os exames realizados por laboratórios e laudos realizados por médicos, para diminuir o tempo entre diagnóstico de câncer de mama e início do tratamento. Essa IA utiliza essa inteligência justamente para otimizar o tratamento. Inclusive, já foram publicados estudos na *National Library of Medicine* associando a demora no tratamento entre o diagnóstico e o início com a taxa de mortalidade muito alta. Então, nesse caso a gente não teria uma IA intervindo diretamente no diagnóstico, por exemplo, ou no exame, mas, de fato, a gente teria um benefício direto dessa inteligência artificial, a um risco zero.

A gente vai navegar por esse exemplo aqui hoje para entender um pouco por que o texto do PL 2.338, na forma como proposto, precisa ainda de alguns ajustes para que possamos incentivar que tecnologias como essa, que salvam vidas, sejam desenvolvidas aqui em território nacional, e não permitir que mais talentos como esse fujam do Brasil à procura de um ambiente jurídico mais propício para desenvolvê-las, como os Estados Unidos, por exemplo, que têm muitos incentivos, tanto em nível privado quanto público, para desenvolvimento dessas novas tecnologias.

O cerne aqui da nossa preocupação em relação a esse PL está em três pontos principais, que, na medida possível do tempo, vou tentar abordar aqui. O primeiro é a falta de compreensão sobre ciclo de vida da IA e definição dos agentes; o segundo é a categorização de risco, muito bem explicada pela Profa. Dra. Ana; o terceiro é a lógica da responsabilidade civil aplicada no PL.

Em relação às definições, em primeiro lugar, o PL traz apenas duas figuras como agentes da IA: o fornecedor e o operador. De antemão a gente já tem que a nomenclatura do operador já está presente na LGPD e, com isso, a gente pode ter uma confusão muito grande no momento de incidência de ambas as leis, tanto a LGPD quanto a lei que agora a gente discute, numa situação específica, porque o operador de dados lá pode não ser o operador da IA aqui e vice-versa. Então, só por isso a definição e as nomenclaturas já mereciam ser revistas.

Em segundo lugar, há uma dificuldade ainda muito grande de cravarmos uma definição para os agentes. Essa discussão, inclusive em nível internacional, ainda está muito quente, mas pelo menos uma das figuras merece ser destacada desse bolo, que é a figura justamente do desenvolvedor da IA. Muitas vezes esse desenvolvedor não faz parte mais da cadeia de distribuição e de uso da inteligência artificial,



SENADO FEDERAL
Secretaria-Geral da Mesa

mas, pela lógica proposta no texto legislativo, ele participa integralmente da responsabilidade por um dano dentro da cadeia. A gente tem que... De repente, foi apenas uma mente brilhante que desenvolveu uma tecnologia, colocou a mercado, mas ela continua participando daquela cadeia e sendo responsável em caso de uma conduta danosa.

Tem um exemplo que eu gosto bastante de trazer, que é bem elucidativo, que é o exemplo de alguém que, com uma faca, lesiona uma pessoa na rua ou pratica um homicídio, por exemplo. Em nenhuma circunstância do direito brasileiro, as facas Tramontina vão ser responsabilizadas por esse delito, por esse crime ou por essa lesão, porque, da mesma forma como a IA, a gente está falando de uma ferramenta que pode ser utilizada tanto para o bem como para o mal em diferentes circunstâncias. Então, a gente tem que responsabilizar realmente a pessoa, a entidade que causou aquele dano, que foi responsável pelo ato ilícito.

Quando a gente fala da categorização de riscos, aí eu faço coro com a Dra. Ana, porque ficou realmente muito ampla a questão do risco alto. Ele traz um rol, o PL, extensíssimo de atividades consideradas de risco alto, incluindo setores inteiros, como saúde, educação e administração da Justiça. Então, mesmo aquelas IAs que não oferecem risco nenhum, como aquela que eu citei no começo, justamente a que lê os laudos e otimiza o tratamento de câncer de mama, elas são consideradas de alto risco apenas por estarem inseridas dentro do segmento da saúde. Como exemplo, além dessa IA, tem diversas outras. Tem, por exemplo, uma *startup* que desenvolveu um robô que percorre lavouras inteiras identificando pragas específicas para ali poder aplicar o pesticida específico numa área diminuta. A IA é tão inteligente que, ao mesmo tempo em que economiza custos do produtor rural, diminui muito o impacto ambiental. Nesse cenário, ela poderia ser considerada um veículo autônomo, já que ela percorre lavouras, e poderia inclusive ser incluída no rol de risco alto.

Então, o que falta aqui, na nossa percepção, é justamente uma análise mais cuidadosa, mais casuística na atribuição do risco, principalmente porque a tecnologia, como vimos debatendo aqui extensivamente, avança e se transforma de forma constante. E a gente entende que esse tipo de atribuição deveria ser feito de forma setorial por quem realmente entende dos impactos da tecnologia numa atividade específica.

E aí agora, já entrando um pouco no tema da responsabilidade civil, da lógica que foi colocada nesse PL, a gente viu que o texto proposto foi extremamente rígido com essas IAs categorizadas como alto risco, que são tantas como a gente já viu aqui. A gente vê no ordenamento jurídico brasileiro, como a gente já reforçou aqui em diversas exposições, que existem diversos remédios para reparar danos existentes e eficientes já no nosso regramento: nós temos o Código Civil; temos a Constituição Federal; temos a LGPD, que já é bastante protetiva para os titulares de dados pessoais; nós temos o CDC, que já aplica a responsabilidade objetiva e a inversão do ônus da prova no contexto das relações de consumo e do agente hipossuficiente, mesmo tendo IA. Então, se tiver aplicação de IA numa relação de consumo, a gente vai ter já aplicação da responsabilidade objetiva e da inversão do ônus da prova. Por que agora – a gente deixa esta reflexão – todas as relações envolvendo IA mereceriam sair da regra geral de uma responsabilidade subjetiva e seguir a lógica mais consumerista e mais dura de atribuição das responsabilidades? Só porque é uma solução de IA? Mesmo quando a relação não envolve um hipossuficiente ou um consumidor, quando a gente está falando de relações B2B, entre duas empresas, que podem ter a liberdade de negociar contratos, violações, indenizações para o caso específico, vale a pena a gente ter uma lógica tão protecionista que poderia prejudicar o fornecedor da IA, que, por muitas vezes, vai estar no polo mais frágil do ponto de vista econômico do que quem está tomando essa IA? Fica uma lógica muito deturpada em relação às responsabilidades.



SENADO FEDERAL
Secretaria-Geral da Mesa

O texto do PL 2.338 aplica essa dinâmica, que é muito custosa e prejudicial para as nossas pequenas e médias empresas, prevendo essa responsabilidade objetiva e a inversão do ônus da prova em grande parte dos casos. Então, vale mais uma vez frisar que, nesse rol, incluem-se centenas de *startups* que estão desenvolvendo, em solo brasileiro, soluções de ponta e que podem colocar o Brasil no topo da rota internacional de IA, trazendo riqueza para o nosso país. O importante é a gente exportar essa tecnologia. Não teria lógica nenhuma nós importarmos e pagarmos caro por uma tecnologia desenvolvida por brasileiros no exterior. Seria revoltante, inclusive.

(Soa a campanha.)

A SRA. ADRIANA ROLLO – Então, já caminhando para a conclusão, o resumo da nossa contribuição é no sentido de incentivar o viés positivo da regulação da IA, focando os nossos esforços em educação, em governança, em ética *by design*, em transparência, em centralidade humana, que foi bastante falada aqui, em dignidade humana, que foi bastante falada aqui, em não discriminação, em avaliações de impacto algorítmico, enfim, numa abordagem mais orientada a princípios e mais flexível para a indústria nacional poder se desenvolver nesse segmento.

Enquanto nós ainda estamos no início dessa era da IA, com diversas soluções ainda sendo criadas, nós não estamos nem perto de entender o real potencial e, principalmente, o real risco dessa tecnologia. Por isso é que nós devemos nos prevenir com uma regulação que seja eficiente e que permita e incentive o desenvolvimento de novas ferramentas de IA dentro de um sistema jurídico e tecnológico...

(Soa a campanha.)

A SRA. ADRIANA ROLLO – ... que seja ético e seguro, centralizado na pessoa humana.

Infelizmente, nós não podemos, como disse o Senador Marcos Pontes muito bem, antecipar e controlar todos os riscos possíveis ou impossíveis que a IA pode trazer, mas o fato é que, quanto mais nós tentarmos controlar e antecipar esses riscos que nós desconhecemos ainda num texto legal, mais nós vamos correr o risco real de engessar a inovação, aumentando a chance de ficarmos de fora nessa corrida internacional pela IA.

Muito obrigada, Senador.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Muito obrigado, Dra. Adriana. Parabéns pela apresentação.

Realmente alguns pontos me chamaram a atenção. Essa necessidade de alinhamento das leis já tem aparecido em várias dessas nossas audiências, por exemplo, com a Lei Geral de Proteção de Dados, que já existe e já está em funcionamento. Então, não faria sentido repetir ou, pior ainda, entrar em convergência nem em choque com o resultado. As responsabilidades também são extremamente importantes – eu gostei do exemplo da faca, que é uma ferramenta afinal de contas – e também a questão do risco, da categorização de risco alto ali.

Sempre o que me chamou a atenção, na primeira vez que eu olhei, foi a parte de educação. É lógico que a educação tem uma importância grande. A impressão que eu tive é que aquilo ali foi colocado, setores que são importantes – e o setor da educação é importante –, mas a aplicação, o resultado dessa aplicação é que tem um risco envolvido. E toda essa categorização do risco, a identificação, a priorização, tudo isso que é feito, que nós fazemos, como uma ação que precisa ser feita com relação a risco para mitigar o problema – a questão de que você aceita o risco, passa para um terceiro etc., ou elimina aquela possibilidade –, isso aí tudo tem que ser levado em conta. Realmente, essa análise técnica eu acho que era o que realmente precisávamos trabalhar dentro dessa lei.



SENADO FEDERAL
Secretaria-Geral da Mesa

E há também a questão do investimento. Sempre se brigou muito com relação a investimento para pesquisa e desenvolvimento no país. Eu já falei isto tantas vezes, mas não se consegue desenvolver um país se não tiver o investimento adequado em áreas estratégicas, que são educação, ciência, tecnologia, inovação. Simplesmente é impossível! É só olhar, nos países desenvolvidos hoje, o que eles têm em comum: todos eles têm em comum exatamente o foco nessas áreas. Então, esse cuidado com o investimento, com a proteção do desenvolvimento das empresas do setor é muito importante.

Aliás, há uma boa notícia em relação a isso aqui. Não na semana que vem, mas na outra, a gente já tem lá na Comissão de Constituição e Justiça a apreciação e, eu espero, a aprovação unânime – eu acho que ninguém é contra pesquisa e desenvolvimento, nenhum partido – da PEC 31, a chamada PEC da ciência, para que o Governo, todas as instituições, todos os entes federais, todos os entes no país, os públicos e as entidades privadas, pensem em desenvolvimento de pesquisa e desenvolvimento dos investimentos. Isso, para que o Brasil saia do patamar ali de 1,2%, 1,14% de investimento em pesquisa e desenvolvimento no país, somando investimento público e investimento privado – é bom ressaltar isso, porque as pessoas confundem –, para chegar próximo ao patamar da OCDE, em que a média dos países é de 2,71% do PIB. O Brasil está com 1,14% do PIB, lá embaixo. Em dez anos... O que essa PEC propõe é que o Governo e todas as entidades façam um esforço de articulação junto ao setor privado, etc., para que esse investimento chegue ao nível de 2,5% no mínimo em dez anos, o que nos levaria aí a um nível, a um patamar melhor, o que vai favorecer todas as áreas em todos os setores que utilizam e vão utilizar cada vez mais tecnologia. É ainda distante de uma Coreia do Sul, com 5% do PIB, ou de Israel, com 5% do PIB, mas já mais próximo da OCDE, o que nos traz aí expectativas boas. Eu espero que essa PEC seja aprovada rapidamente para que a gente comece a andar nesse sentido, e finalmente com pragmatismo.

Bom, nós temos mais um debatedor antes de a gente entrar na fase das perguntas, que é o Carlos Affonso. O Sr. Carlos Affonso de Souza é Consultor da Associação Brasileira de Internet (Abranet).

Carlos, obrigado pela paciência, aguardando o tempo todo – o que é importante, porque vai pegando uma série de ideias também – para os seus dez minutos. Por favor, controle o tempo aí remotamente. Obrigado.

O SR. CARLOS AFFONSO DE SOUZA (Para expor. *Por videoconferência.*) – Obrigado, Senador Astronauta Marcos Pontes.

Do lado de cá, devo dizer que o fato de ficar para este momento da audiência também tem as suas vantagens, como o Senador bem lembrou. Tive a oportunidade de aprender muito com os colegas, com os amigos todos que puderam falar antes, aqui nessa audiência.

Queria só estender aqui os agradecimentos ao Senador Eduardo Gomes pelo brilhante trabalho que, junto com o Senador Astronauta Marcos Pontes, tem feito nessa Comissão Temporária de Inteligência Artificial.

E quero dizer que hoje represento aqui a Abranet (Associação Brasileira de Internet). A Abranet é uma associação criada em 1996 – de certa maneira, a história da Abranet se confunde com a história da internet brasileira –, contando hoje com mais de 400 associados nos mais diferentes campos de atuação ligados à internet e à tecnologia, de *fintech*, *paytech*, provedores de acesso, provedores de conteúdo das mais diferentes plataformas digitais hoje na rede.

E é justamente esse conjunto de atividades bastante amplo da associação que permite, de certa maneira, fazer com que a associação possa ter um olhar sobre esse tema que é um olhar que gostaríamos de comunicar muito claramente na largada – e sei que esse tema tem aparecido nas mais diferentes audiências dessa Comissão, mas é importante deixar isto claro –: inteligência artificial não é



SENADO FEDERAL
Secretaria-Geral da Mesa

um setor. Inteligência artificial é, como o próprio Senador muito bem colocou, um tema transversal. É um tema que vai impactar as entidades, em especial no setor privado, as empresas de logística, de manutenção, de produção, de conteúdo, de varejo.

De certa maneira, dizer que uma empresa é uma empresa de inteligência artificial talvez seja uma frase que, em pouco tempo, não vai fazer muito sentido. Toda empresa vai ser uma empresa de inteligência artificial. De certa maneira, me parece que, num curto espaço de tempo – assim como hoje nos parece muito estranho uma empresa que não esteja na internet –, uma empresa que não aplique, de maneira importante para as suas atividades, determinadas soluções de inteligência artificial será objeto de enorme espanto. Então, de certa forma, acho importante comunicar esse ponto de que nós estamos falando num futuro no qual inteligência artificial não será a exceção; inteligência artificial será a regra. E, justamente por isso, essa visão precisa estar plasmada em todas as atividades, em todas as iniciativas que saem por parte do poder público e que instruem a própria iniciativa, o próprio pensar da regulação como um todo.

E, justamente fazendo esse olhar para a inteligência artificial como não sendo a exceção, como sendo a regra, eu gostaria de já puxar alguns pontos. E aí, procurando fugir um pouco de temas que já foram explorados aqui nessa audiência hoje, eu vou acabar me dedicando mais a alguns pontos em especial do PL 2.338.

E aqui, desde já parablenizo o trabalho feito pela Comissão de Juristas, capitaneado pelo Ministro Cueva, que entregou aqui um relatório e uma proposta de texto que certamente leva adiante essa discussão sobre regulação no Brasil.

Mas me parece que é preciso, enfim, como acontece com todo processo legislativo aqui, se debruçar sobre alguns pontos de atenção do texto. E eu chamaria a atenção, logo na largada, no projeto, sobre os personagens desse projeto de lei. E o primeiro deles é a chamada pessoa afetada por sistemas de inteligência artificial.

Parece-me que aqui nós temos, na própria maneira pela qual o PL se refere a esse personagem, a esse protagonista, no final das contas, do PL 2.338, algo que mereceria um olhar cuidadoso, porque, quando se fala em "pessoa afetada por sistemas de inteligência artificial", parece, de novo, que estamos falando de algo excepcional, de pessoas que são afetadas por inteligência artificial como sendo a exceção, e sabemos que todos nós seremos, entre aspas, "afetados" por inteligência artificial diversas vezes, no nosso dia a dia; IA será rotina, não será excepcional. Então, por isso, na busca desse protagonista da legislação, me parece que a opção por "pessoa afetada por sistema de inteligência artificial" não seria a melhor solução, me parece que nós cravaríamos no texto da lei uma redação que tem caráter excepcional.

De certa maneira – e aqui sabendo que muito do pensamento sobre o PL 2.338 deriva de uma matriz próxima do que se pensou para a Lei Geral de Proteção de Dados –, na Lei Geral de Proteção de Dados, temos o titular de dados pessoais. E "o titular de dados pessoais" é uma expressão muito mais neutra e que, de certa maneira, não carrega esse tipo de compreensão que me parece que nós gostaríamos de evitar numa eventual solução regulatória tirada a partir do PL 2.338.

Ainda seguindo um olhar sobre personagens do PL 2.338 ou de qualquer solução regulatória sobre inteligência artificial, me parece que, quando falamos sobre as figuras dos agentes de sistemas de inteligência artificial, na escolha das figuras do fornecedor e do operador de sistemas de inteligência artificial – e esse ponto já apareceu numa intervenção anterior –, essas duas figuras parecem ser insuficientes para trabalhar com a complexidade de atores que vão estar inseridos na linha de



SENADO FEDERAL
Secretaria-Geral da Mesa

desenvolvimento, treinamento, customização, colocação no mercado, manutenção de aplicações relacionadas à inteligência artificial. E esse é um outro ponto que me parece ser um ponto de atenção.

De certa maneira, na LGPD, a figura de controlador e operador do tratamento de dados pessoais, como duas espécies do gênero "agentes de tratamento de dados", é uma estrutura que foi bem recepcionada e bem desenvolvida por parte das autoridades e do mercado, das empresas, das entidades que tratam dados pessoais. No que diz respeito a sistemas de inteligência artificial, me parece que nós precisamos olhar com mais cuidado, em especial se nós pensarmos nas entidades que desenvolvem um modelo fundacional em cima do qual eu vou desenvolver uma série de treinamentos e customizações que vão estar embarcadas em diferentes aplicações, porque serão essas aplicações que chegarão ao mercado. Então, de certa maneira, entender qual é o papel de uma entidade que desenvolve um modelo fundacional – e aqui a gente tem as mais diferentes situações que aparecem já neste momento, já neste instante na questão de desenvolvimento de inteligência artificial, com o Stable Diffusion, com o GPT, com o Llama... Então, acho que é importante a gente pensar qual é o papel desses modelos fundacionais que vão servir para a criação de um sem número de aplicações que vão ser customizadas e treinadas das mais diferentes formas e de que maneira as figuras de fornecedor e operador que nós temos hoje no PL 2.338 açambram essas situações como um todo.

Já caminhando aqui para o finalzinho da minha fala, até para não pressionar muito o tempo de todos, parece-me – e essa é uma fala que vem muito claramente dos posicionamentos mais recentes da Abranet (Associação Brasileira de Internet) – que é necessário haver aqui uma maior coordenação sobre as diferentes atividades envolvendo inteligência artificial nos diferentes Poderes. Nós temos aqui no Poder Legislativo o debate feito nesta Comissão de uma maneira muito destacada. Nós temos, no Poder Executivo, a divulgação da Estratégia Brasileira de Inteligência Artificial e, em diferentes segmentos, mas olhando aqui em especial a ANPD, consulta pública sobre *sandboxes* como uma iniciativa também muito positiva nessa área. E o Judiciário, numa manifestação muito recente do Ministro Barroso, na Presidência do Supremo Tribunal Federal, falou sobre a necessidade de o Poder Judiciário contar com ferramentas que possam tornar o dia a dia do magistrado e da magistrada mais eficiente.

Então, de certa forma, nós precisamos aqui ter uma coordenação para que o Brasil tenha uma visão que lhe seja própria de inteligência artificial e que possa explorar a identidade do país nesse sentido. Parece-me que a estratégia brasileira deu um primeiro passo, mas, em especial, pensando não só na discussão regulatória, mas no posicionamento do país frente às negociações em um cenário internacional no qual inteligência artificial é cada vez mais um tema crucial, parece-me – e olhando em especial para o G20, que será sediado no Rio de Janeiro em setembro do ano que vem – que nós precisamos realmente olhar com bastante cuidado o DNA brasileiro, pensando em sustentabilidade, diversidade e a experiência multissetorial que o Brasil tem no diálogo sobre novas tecnologias.

Com isso, encerro aqui, Senador, mais uma vez agradecendo pela oportunidade de participar desse debate, saudando a Comissão Temporária de Inteligência Artificial por esse importante trabalho que faz para que a gente possa ter um diálogo cada vez mais aperfeiçoado, cada vez mais diverso sobre esse tema tão relevante.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Obrigado, Sr. Carlos Affonso de Souza, Consultor da Associação Brasileira de Internet (Abranet).

Bom, terminadas as apresentações dos nossos debatedores, eu gostaria de agora apresentar... Acho que já receberam também as perguntas do nosso pessoal que participou, mas eu vou fazer a leitura dessas perguntas, mesmo porque os nossos debatedores no remoto não têm as perguntas em mão. Então, eu vou fazer a leitura dessas perguntas.



SENADO FEDERAL
Secretaria-Geral da Mesa

E, dado o nosso tempo, como a gente vai fazer essa parte final aqui? Eu vou ler as perguntas, e os debatedores olhem quais as que vocês queiram destacar. Depois, eu vou passar para as considerações finais para cada debatedor e, nas considerações finais, então, eu peço que façam a abordagem dessas questões. Pode ser assim?

Então, vamos às perguntas que nós recebemos através do e-Cidadania, através do Portal do Senado e também do telefone de contato. Aliás, eu quero agradecer às pessoas que participaram.

Daniel Corrêa, de Minas Gerais: "No passado a regulamentação da profissão de [...] [tecnologia da informação] foi negada para não limitar a área, regulamentar a IA não trará limitações [...]?". Então, isso aí esteve nos nossos debates aqui.

Lizandro Mello, do Rio Grande do Sul: "[...] o art. 7, inciso XXVII, da Constituição (proteção do trabalhador em face da automação) [é central]. Como inseri-lo nas leis sobre [...] [inteligência artificial]?"

Ângelo Gabriel, do Rio Grande do Sul também: "Quais os benefícios e malefícios que as inteligências artificiais podem trazer ao meio ambiente?"

Edissa Lília, de São Paulo: "Existem planos para proteger artistas e seu direito intelectual, já que as [...] [inteligências artificiais] usam as artes em seu banco de dados sem permissão?"

Rodolfo da Silva, de São Paulo: "Como garantir eficácia e segurança na utilização de [...] [inteligência artificial] em sistemas de alto risco, como de saúde e bancário, por exemplo?"

Igor Lázaro, do Paraná: "Até que ponto as tecnologias em [...] [inteligência artificial] podem ser incluídas na estrutura do estado [nos dias de] hoje [...]?"

Ícaro Ferreira, do Ceará: "Quais são as medidas de prevenção [que] o governo está adotando para [...] [impedir] os crimes virtuais [...] [que usam inteligências artificiais] para gerar, voz e vídeo?"

Yasmin Amparo, do Rio de Janeiro: "Como a regulação brasileira vai lidar com o uso de [...] [inteligência artificial] que poderá vir de fora do nosso país?"

Vinícius Pereira, da Bahia: "Quais são os riscos que as [...] [inteligências artificiais] oferecem às eleições? Que medidas devem ser tomadas para regular as propagandas políticas utilizando [...] [inteligência artificial]?"

Richardson Silva, de Minas Gerais: "Quais são os princípios fundamentais que devem guiar a regulação da inteligência artificial?"

Emanuel Polari, da Paraíba: "Qual é o equilíbrio ideal entre o avanço da inteligência artificial e a preservação dos valores éticos, da privacidade e da responsabilidade?"

Sandro Júnior, do Rio Grande do Sul: "Quais serão os órgãos responsáveis por acompanharem o desenvolvimento da inteligência artificial?"

Fábio Lima, de São Paulo: "Como vamos garantir a integridade econômica do país, caso a [...] [inteligência artificial] [...] [substitua] empregos de baixa renda?"

Essas são as perguntas.

E, agora, eu vou passar a palavra para cada um dos nossos debatedores para suas considerações finais. Provavelmente a ordem não vai ser a mesma, porque a ordem foi um tanto modificada desde o começo. Então, eu vou fazer aqui na sequência que nós temos aqui no presencial e, depois, lá no remoto para as considerações finais.

E eu vou começar, então, com a Fernanda Rodrigues, que é Coordenadora de Pesquisa do Instituto de Referência em Internet e Sociedade, para suas considerações finais e comentários para alguma pergunta que queira comentar.

A SRA. FERNANDA RODRIGUES – Quanto tempo tenho?



SENADO FEDERAL
Secretaria-Geral da Mesa

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Dois minutos.

A SRA. FERNANDA RODRIGUES (Para expor.) – Na verdade, eu queria só aproveitar para agradecer novamente pela oportunidade.

Eu fico realmente muito feliz de estar do lado de representantes de diferentes setores aqui na mesa, porque, assim como a Cynthia, eu acredito realmente que a governança multissetorial pode ser o caminho mais acertado para resolver os problemas relacionados à IA – e não só multissetorial, mas multidisciplinar, interdisciplinar também, porque, como a gente já falou várias vezes aqui, a tecnologia é transversal. Então, é importante, sim, que pessoas da área técnica estejam à frente deste debate – a Adriana trouxe isso muito bem –, porque realmente são as pessoas que entendem a fundo isto, entendem a fundo como a tecnologia funciona; e também, na medida em que essa tecnologia vai refletir e impactar diretamente na sociedade, que este debate seja tanto técnico quanto social, também considerando aspectos éticos e sociais. E aí a gente precisa trazer, portanto, juristas, cientistas sociais, antropólogos, todas essas áreas, para que a gente possa conversar e pensar uma tecnologia que parta...

(Soa a campanha.)

A SRA. FERNANDA RODRIGUES – ... destes dois caminhos: a interdisciplinaridade e a multissetorialidade.

Eu também estou de pleno acordo. Queria finalizar dizendo que eu estou de pleno acordo de que a tecnologia avança realmente muito rápido e que o seu futuro é incerto, a gente não tem como saber o que vai ser desenvolvido na sequência. Inclusive eu quero ter oportunidade de conversar mais com o setor técnico e com o setor privado a respeito dessas possibilidades futuras, mas isso também não significa fechar os olhos porque a gente já tem visto sobre a inteligência artificial atualmente. Então, para além de um exercício preditivo, um exercício talvez realístico de olhar para o que a gente já tem e também precisa ser adequadamente endereçado.

Eu acredito que é, sim, possível construir, portanto, uma regulação que seja firme nos direitos e responsabilidades, ao mesmo tempo que seja flexível para permitir uma atualização conforme a tecnologia vai avançando.

Obrigada.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Muito obrigado. A Fernanda Rodrigues é Coordenadora de pesquisa e Pesquisadora do Instituto de Referência em Internet e Sociedade (Iris). Obrigada pela participação.

Eu passo, então, a palavra para as considerações finais ao Dr. Rodrigo Badaró, Conselheiro Nacional do Ministério Público.

O SR. RODRIGO BADARÓ (Para expor.) – Presidente, mais uma vez agradeço a todos, foi muito enriquecedor este debate, estou muito feliz em estar aqui. Queria até aproveitar esta oportunidade final de até dar um *spoiler* e contar rapidamente a situação. Eu sou relator, Senador, de um pedido de providência de um advogado com um pedido de liminar que tentou impedir o uso de inteligência artificial por todo o Ministério Público brasileiro.

Eu ouvi todo o Ministério Público brasileiro junto com o Conselheiro Moacyr Rey e, na próxima sessão do dia 14 de novembro, do nosso Conselho Nacional do Ministério Público, que é o órgão máximo do Ministério Público brasileiro, nós iremos apresentar uma recomendação, não é resolução, exatamente numa linha subjetiva, principiológica, para tentar dar um mínimo norte, nesse caso setorial, para o Ministério Público brasileiro.



SENADO FEDERAL
Secretaria-Geral da Mesa

Estamos em debate avançado também com nosso co-órgão irmão, que é o Conselho Nacional de Justiça, com o grande Conselheiro, inclusive indicado por esta Casa, o Bandeira de Mello, para eventualmente criar um grupo de trabalho muito em breve. Espero que ele consolide esse projeto junto ao CNJ, entre o CNJ e o CNMP, que haverá certamente um diálogo interinstitucional e, dentro de uma regulação interna setorial, obviamente respeitando qualquer lei federal que venha a ser aprovado.

Com essas considerações, tentando ajudar V. Exa., Senador, com alguma resposta...

(Soa a campanha.)

O SR. RODRIGO BADARÓ – ... o Richardson pergunta aqui de Minas Gerais, da minha terra natal: "Quais são os princípios fundamentais que devem guiar a regulação da Inteligência Artificial?" Eu acho difícil criar princípios, acho que a gente já tem – como eu coloquei na minha exposição – acho que são os princípios já colocados na Constituição Federal e desenvolvidos, de matiz constitucional, até como proteção de dados, que virou direito fundamental da dignidade humana, da transparência, da liberdade da atividade econômica, da liberdade do direito à proteção à privacidade. Todos esses já estão inseridos na Constituição. Obviamente pode-se discutir a interpretação desses princípios com questões já colocadas, por exemplo, como colocada pela Dra. Estela.

E tentando responder mais uma pergunta, a do Sandro: "Quais serão os órgãos responsáveis por acompanhar o desenvolvimento da inteligência artificial?" Acho que essa resposta eu não tenho – acho que ninguém tem ainda. Acho que já há um trabalho e investimento da Anatel, uma consolidação da Anatel; há também um trabalho e investimento da NPD; e há um trabalho hercúleo do Ministério da Justiça também, e do próprio Senado, como órgão parlamentar, no debate desse tema, mas a gente não tem ainda um órgão de desenvolvimento, muito menos ainda um órgão regulatório que faz parte, inclusive, do debate que V. Exa. promove hoje.

Muito obrigado, Senador.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Obrigado, obrigado, Dr. Rodrigo Badaró, Conselheiro Nacional do Ministério Público.

Eu passo, na sequência, para as suas considerações finais, à Dra. Adriana Rollo, Líder da Comissão Especial de Regulação de Inteligência Artificial da Associação Internacional de Inteligência Artificial.

A SRA. ADRIANA ROLLO (Para expor.) – Obrigada, Senador.

Eu vou pegar carona aqui numa pergunta muito pertinente do Emanuel Polari, da Paraíba, que pergunta: "Qual é o equilíbrio ideal entre o avanço da Inteligência Artificial e a preservação dos valores éticos, da privacidade e da responsabilidade?" Bom, equilíbrio ideal a gente nunca vai conseguir em nenhum momento, mas é importante aqui a gente sopesar direitos e garantias fundamentais dos indivíduos e realmente um avanço tecnológico que pode impactar positivamente milhares de outros. Então, por isso que a gente tem aqui diversos princípios já aplicáveis ao uso de IA, como bem falou o Dr. Rodrigo. Então, a gente tem a preocupação aqui do Emanuel sobre privacidade. Está endereçada em diversas leis. A gente tem a Lei Carolina Dieckmann, a gente tem a LGPD também, tem os princípios fundamentais da Constituição. Então, nós já temos esses mecanismos e remédios para poder aplicar ao tipo de dano que ocorra a uma pessoa física. Por isso que a gente tem que ter um pouco de...

(Soa a campanha.)

A SRA. ADRIANA ROLLO – ... consciência e direcionamento para que a inteligência artificial também possa ser responsável por impactar positivamente diversas pessoas dentro de um regime de responsabilidade que seja coerente e não punitivo e prescritivo.



SENADO FEDERAL
Secretaria-Geral da Mesa

Obrigada, Senador. Foi um prazer debater com vocês aqui. Esta manhã foi muito produtiva.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Muito obrigado, Dra. Adriana Rollo, Líder da Comissão Especial de Regulação de Inteligência Artificial da Associação Internacional de Inteligência Artificial.

Na sequência, para as suas considerações finais, eu passo a palavra para a Profa. Dra. Ana Carla Bliacheriene, da Universidade de São Paulo.

A SRA. ANA CARLA BLIACHERIENE (Para expor.) – Obrigada. Eu queria agradecer a oportunidade de estar aqui nesta discussão com colegas que nos precederam em outros dias, com todos que construíram esse cabedal. Hoje, se a gente pode discutir nos termos que estamos discutindo, é porque houve outros que construíram essa base para a discussão. Então, quero agradecer por estar aqui nesta mesa hoje; agradecer a esta Casa; ao Senador Eduardo Gomes; e a V. Exa. pelo convite e oportunidade que dá à Universidade de São Paulo.

Quero aproveitar aqui duas perguntas.

Uma é da Edissa Lília, que pergunta em relação à proteção do direito da propriedade intelectual dos artistas. O meu entendimento é que esse tema não pode ser tratado numa lei geral de regulação de IA; precisa ser uma lei específica e com uma rodada de discussões no nível do que está sendo esta, porque é um tema árduo.

Outra é do Rodolfo, que fala: "Como garantir [...] segurança na utilização de IA em sistemas de alto risco, como de saúde e bancário [...]?". Talvez, Rodolfo...

(Soa a campainha.)

A SRA. ANA CARLA BLIACHERIENE – ... a gente pensar que a questão não é a área que é de alto risco; é a atividade. Não é a área que tem alto risco. É perigoso andar de avião, é perigoso andar de carro. Muitas pessoas no Brasil morrem por dia por acidentes de trânsito. Não é a área, mas é a atividade dentro da área. Isso só se faz de um mapeamento adequado de riscos.

E quero dizer que o meu desejo realmente é que essa lei possa cobrir o impacto da IA no setor público, que é uma área em que tenho uma especial atuação. Eu tenho atuado muito fortemente com tribunais de justiça e com tribunais de contas. Enxergo, Excelência, um campo espetacular de boa prestação de serviço público, através do uso massivo de IA.

Então, reconhecendo a centralidade humana, reconhecendo o humanismo e reconhecendo a necessidade de cuidar do povo brasileiro e das pessoas do Brasil, a IA não é um inimigo da política pública; muito pelo contrário: a IA será um grande parceiro de políticas públicas efetivas e que cuida das pessoas.

Nós precisamos disso.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Muito obrigado. Muito obrigado e parabéns pelos comentários.

A Dra. Ana Carla Bliacheriene é Professora da Universidade de São Paulo.

Na sequência, eu passo a palavra à Dra. Marcela Mattiuzzo, Conselheira do Instituto Brasileiro de Estudos de Concorrência, Consumo e Comércio Internacional.

Por favor.

A SRA. MARCELA MATTIUZZO (Para expor.) – Muitíssimo obrigada, Senador.

Enfim, novamente agradecendo a oportunidade para o Ibrac participar desta discussão, eu queria endereçar algumas perguntas – a gente tem algumas, várias, na verdade – sobre inovação e regulação



SENADO FEDERAL
Secretaria-Geral da Mesa

e como fazer alguma proposta regulatória pode levar a um problema do ponto de vista da inovação. Daniel Corrêa fala sobre isso, Emanuel Polari também, enfim.

O Ibrac, de fato, entende que isso é muito relevante, até por isso que a gente enfatizou tanto na nossa fala anterior a importância de a discussão regulatória andar lado a lado com a discussão do desenvolvimento tecnológico. Então, a gente entende que isso é primordial, mas também acho que é importante ressaltar que é um pouco falacioso, talvez seja um pouco simplista a gente meramente achar que tem uma incompatibilidade completa entre essas duas coisas.

Enfim, tem muitos exemplos simples de regulações que levaram em algum sentido a alguma restrição inovativa que a gente considera que são bons. Por exemplo, um cinto de segurança no carro.

(Soa a campainha.)

A SRA. MARCELA MATTIUZZO – Com certeza isso leva a restrições de inovação e implica custos, inclusive para um fabricante, mas hoje se entende que isso faz sentido.

Então, a dificuldade que a gente tem aqui é como balancear essas duas coisas e como a gente encontra uma regulação que não leve a, enfim, problemas intransponíveis do lado da inovação.

Muito obrigada novamente.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Obrigado, Dra. Marcela Mattiuze, Conselheira do Instituto Brasileiro de Estudos de Concorrência, Consumo e Comércio Internacional (Ibrac).

Na sequência, eu passo a palavra à Dra. Estela Aranha, Assessora Especial de Direitos Digitais do Ministério da Justiça e Segurança Pública.

Por favor.

A SRA. ESTELA ARANHA (Para expor.) – Obrigada, Senador.

Então, vou falar um pouco também sobre a segurança e eficácia dos sistemas de alto risco, porque realmente existem sistemas de alto risco quando o uso de uma ferramenta pode trazer riscos de vida, integridade física ou aos direitos fundamentais, por isso que a gente fala na garantia do *trustworthy*. A IA tem que entregar o que ela promete com segurança, o resultado tem que ser... Se você faz um exame com a IA, o resultado tem que ser adequado. É falsa a ideia de que ninguém garante esses riscos no mundo inteiro. Eu vou dar um exemplo da saúde, que foi colocada aqui.

Os Estados Unidos têm uma legislação, por exemplo, mais liberal em relação à regulação, mas e a IA e saúde? Se você olhar os manuais do FDA, para você fazer a verificação de um *device*, um dispositivo médico de saúde para ser usado no sistema de saúde, eu vou dar algumas perguntas aqui que são dadas na documentação de gerenciamento de risco que eles pedem dos dados. Qual a garantia da qualidade dos dados? Quais foram usados no *machine learning*?

(Soa a campainha.)

A SRA. ESTELA ARANHA – De onde vieram, bancos de dados, a representatividade, porque tem a coisa da estatística. Como é a documentação? Como é feito o treinamento? Segurança: quais as práticas de segurança cibernética? Se esses resultados são acompanhados, quais ajustes? Quais as mitigações de riscos e *bias*? Avaliação da acurácia, da precisão, de índices.

Quais os métodos... O desempenho, por exemplo: quais são os testes realizados que mostram que o desempenho do dispositivo é igual àquele em condições clinicamente relevantes aos estudos clínicos? Você tem que comparar o que a máquina dá com o que clinicamente se coloca. Enfim, quais os médicos responsáveis? Como são feitos os estudos clínicos? Tem estudos científicos sobre o uso de



SENADO FEDERAL
Secretaria-Geral da Mesa

informações claras e apresentadas para o paciente? A supervisão humana: quando a supervisão humana é obrigatória? Qual foco é colocado no desempenho da equipe de IA e supervisão humana? Quem é a equipe? Quais as etapas de desenvolvimento de aplicação que a verificação humana teve? Como foram feitas essas intervenções e de que forma?

Enfim, todas essas questões de regulação de que a gente está falando aqui, que a gente colocou ou nos direitos ou nos requisitos para que você aprove uma IA de alto risco, elas já existem, já são debatidas no mundo. Ninguém pode colocar uma inteligência artificial que coloque em risco, por exemplo, a saúde humana, o resultado de um exame que depois vai ter impacto ou, obviamente, qualquer robótica que possa pôr em risco a vida, a integridade física ou a questão dos direitos humanos sem avaliar se aquele dispositivo ou aquele algoritmo, independentemente da forma com que ele chegue, está alcançando esse resultado. Isso é o mínimo para que a gente tenha o desenvolvimento que todos queremos de uma IA com qualidade, segura e que ajude no desenvolvimento humano.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Obrigado, Dra. Estela Aranha, Assessora Especial de Direitos Digitais do Ministério da Justiça e Segurança Pública. Muito obrigado.

Agora passamos às considerações finais dos nossos debatedores no remoto, começando com o Leonardo Netto Parentoni, Professor da Universidade Federal de Minas Gerais. Para suas considerações finais, por favor.

O SR. LEONARDO NETTO PARENTONI (Para expor. *Por videoconferência*.) – Primeiro eu agradeço, uma vez mais, a oportunidade de participar desta audiência pública e contribuir. É sempre muito não só prudente como produtivo ouvir bastante antes de decidir. É o que o Senado e vários outros órgãos estão fazendo, e parabênizo-os por isso.

Se eu puder reduzir a uma frase, é: regular o que de fato existe na prática, e não regular uma ideia; regular o mercado que temos hoje e como a tecnologia de fato funciona, efetivamente conhecendo o que ela pode e o que ela não pode entregar, ao invés de descer a minúcias, que talvez fiquem ultrapassadas muito rápido, pois a tecnologia não só é majoritariamente produzida fora do Brasil como muda muito rápido.

E nesse minuto que me falta, parece-me que o PL 2.338 não considerou o ponto central de que o que se regula não é o setor, como já foi dito aqui, mas uma aplicação específica dentro de cada setor. Setores diferentes podem usar o mesmo sistema de IA. Simplesmente por ser setor A ou B, isso não deveria ser classificado como alto risco.

E, por fim, a fala das colegas Ana Carla e Adriana Rollo: responsabilizar um sistema que simplesmente lê um exame de imagem médico não faz qualquer sentido do ponto de vista fático, econômico e jurídico, porque ele não substitui a decisão humana. Ele simplesmente sugere ao médico uma entre várias posturas clínicas. Seria o equivalente a responsabilizar, como foi dito, o fabricante da faca pelo homicídio ou pelo dano. Só faz sentido um maior peso da regulação quando de fato o sistema de IA substitui ou todas as etapas ou majoritariamente, a maior parte das etapas da decisão humana. E o PL não foi construído considerando isso.

Muito obrigado, e boa manhã.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Obrigado ao Prof. Leonardo Parentoni, da Universidade Federal de Minas Gerais.

Eu passo a palavra, então, para suas considerações finais, ao Dr. Carlos Affonso de Souza, Consultor da Associação Brasileira de Internet (Abranet).

Tem a palavra.



SENADO FEDERAL
Secretaria-Geral da Mesa

O SR. CARLOS AFFONSO DE SOUZA (Para expor. *Por videoconferência.*) – Obrigado, Senador.

Eu acho que, nesta consideração final, eu vou buscar uma das perguntas da Edissa Lilia sobre direito autoral, e vou focar nesse tema.

Eu tive a oportunidade de voltar recentemente do Japão, do Fórum de Governança da Internet, global, da ONU, e foi muito interessante perceber como o debate no Japão olha para esse tema de direitos autorais. São dois pontos que eu queria conectar com o debate do nosso PL.

O primeiro é um debate sobre alcançar um certo padrão que possa ser compartilhado na internet como um todo, que permita você identificar determinados conteúdos, o chamado *originator profile*, que tem sido discutido no Japão, em especial, através de um consórcio para se procurar chegar a um padrão que possa ser adotado na internet de maneira global. E isso serviria tanto para discussões sobre combate à desinformação, como também para identificar um conteúdo protegido por direito autoral, que venha a ser trabalhado por ferramentas de inteligência artificial generativa.

E o segundo é um comentário feito pela Ministra de Inovação e Esporte japonesa que comentava sobre o pensamento de inteligência artificial generativa como em duas fases: uma primeira fase, para treinamento de ferramentas e aplicações; e uma segunda, de disponibilização do conteúdo gerado por essas ferramentas, porque na primeira fase eu teria de certa maneira uma exceção ao direito autoral para que eu pudesse utilizar obras para treinar ferramentas. Mas, uma vez que eu tenha uma solução, algo tirado dessa ferramenta, essa nova obra por assim dizer poderia ser um (*Falha no áudio*)... de direito autoral.

Esse é um tema que é um debate global e acontece no Japão com essas soluções. No nosso PL, existe o art. 42 – e queria só fazer essa indicação. O tema sobre direito autoral aparece no PL 2.338, no art. 42, e me parece que é uma discussão que precisamos fazer de maneira mais aprofundada. E, é claro, o Brasil não pode ficar de fora desse debate global.

Agradeço mais uma vez, Senador.

E aqui quis puxar apenas uma das perguntas para que a gente pudesse atender a todos que participam via o e-Cidadania.

Obrigado, mais uma vez.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Obrigado ao Dr. Carlos Affonso de Souza, Coordenador da Associação Brasileira de Internet.

Eu gostaria, neste momento, finalmente, aqui, de agradecer: agradecer a todos os nossos debatedores, apresentadores, aqui presenciais e remotos também; agradecer a todos aqueles que estão aqui, no auditório, conosco, acompanhando este debate; a todos aqueles que nos acompanharam através da TV Senado, através das redes sociais do Senado e também da Rádio Senado, acompanhando ao vivo esta audiência; e também à nossa Mesa, que nos apoia aqui todas as vezes e que faz com que tudo isso aqui funcione e com que tudo siga da maneira como tem que seguir.

Então, não havendo mais nada a se tratar, eu declaro encerrada esta reunião, com o desejo de um ótimo final de semana a todos também.

Obrigado.

(Iniciada às 10 horas e 06 minutos, a reunião é encerrada às 12 horas e 51 minutos.)