



SENADO FEDERAL  
Secretaria-Geral da Mesa

ATA DA 22ª REUNIÃO DA COMISSÃO TEMPORÁRIA INTERNA SOBRE INTELIGÊNCIA ARTIFICIAL NO BRASIL DA 2ª SESSÃO LEGISLATIVA ORDINÁRIA DA 57ª LEGISLATURA, REALIZADA EM 04 DE SETEMBRO DE 2024, QUARTA-FEIRA, NO SENADO FEDERAL, ANEXO II, ALA SENADOR ALEXANDRE COSTA, PLENÁRIO Nº 13.

Às quatorze horas e quarenta e dois minutos do dia quatro de setembro de dois mil e vinte e quatro, no Anexo II, Ala Senador Alexandre Costa, Plenário nº 13, sob as Presidências dos Senadores Astronauta Marcos Pontes e Chico Rodrigues, reúne-se a Comissão Temporária Interna sobre Inteligência Artificial no Brasil com a presença dos Senadores Veneziano Vital do Rêgo, Weverton, Vanderlan Cardoso, Nelsinho Trad, Fabiano Contarato, Eduardo Gomes, Laércio Oliveira, Rodrigo Cunha, Izalci Lucas, Alessandro Vieira, Alan Rick, Angelo Coronel, Flávio Arns e Carlos Portinho, e ainda dos Senadores Augusta Brito, Wilder Moraes, Romário, Professora Dorinha Seabra, Paulo Paim e Zenaide Maia, não membros da comissão. Deixam de comparecer os Senadores Styvenson Valentim, André Amaral e Daniella Ribeiro. Deixa de comparecer o Senador Carlos Viana, conforme REQ 498/2024-CDIR. Havendo número regimental, a reunião é aberta. A presidência submete à Comissão a dispensa da leitura e a aprovação da ata da reunião anterior, que é aprovada. Passa-se à Audiência Pública Interativa, atendendo ao Requerimento nº 2, de 2024-CTIA, de autoria Senador Astronauta Marcos Pontes (PL/SP), com a finalidade de debater o tema autorregulação e boas práticas, para instruir o PL 2338/2023, que dispõe sobre o uso da inteligência artificial no Brasil, e seu último relatório, com a participação de: Christina Aires Correa Lima de Siqueira Dias, Advogada Especialista da Confederação Nacional da Indústria (CNI), representante de Antonio Ricardo Alvarez Alban, Presidente da CNI; Ronaldo Lemos, Diretor do Instituto de Tecnologia e Sociedade do Rio de Janeiro (ITS Rio); Luis Fernando Prado, Advogado e professor convidado de pós-graduação *latu sensu* na Faculdade de Direito de Vitória (FDV); Matthew Reisman, Diretor de Privacidade e Política de Dados no Centro de Liderança em Política de Informação de Washington, D.C.; e Courtney Lang, Vice-presidente de Política, Confiança, Dados e Tecnologia no Conselho da Indústria de Tecnologia da Informação de Washington, D.C.. Nada mais havendo a tratar, encerra-se a reunião às dezesseis horas e sete minutos. Após aprovação, a presente Ata será assinada pelo Senhor Presidente e publicada no Diário do Senado Federal, juntamente com a íntegra das notas taquigráficas.

**Senador Astronauta Marcos Pontes**

Presidente em exercício da Comissão Temporária Interna sobre Inteligência Artificial no Brasil

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

[Portal Multimídia \(senado.leg.br\)](https://portal.multimidia.senado.leg.br)

**O SR. PRESIDENTE** (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP. Fala da Presidência.) – Declaro aberta a 22ª reunião da Comissão Temporária Interna sobre



SENADO FEDERAL  
Secretaria-Geral da Mesa

Inteligência Artificial no Brasil, criada pelo Requerimento nº 722, de 2023, com a finalidade de examinar os projetos concernentes ao relatório final aprovado pela Comissão de Juristas, responsável por subsidiar a elaboração de substitutivo sobre inteligência artificial no Brasil, criada pelo Ato do Presidente do Senado Federal nº 4, de 2022, bem como eventuais novos projetos que disciplinem a matéria.

Antes de iniciarmos os 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.*)

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 objetivo de debater o tema autorregulação e boas práticas, para instruir o PL 2.338, de 2023, que dispõe sobre o uso de inteligência artificial no Brasil, e seu último relatório, em cumprimento ao Requerimento nº 2, de 2024, da Comissão Temporária de Inteligência Artificial, de minha autoria.

O público interessado em participar desta audiência pública poderá enviar perguntas ou comentários pelo Portal do Senado, [www.senado.leg.br/ecidadania](http://www.senado.leg.br/ecidadania), tudo junto, ecidadania, ou ligar para 0800 0612211, de novo, 0800 0612211.

Encontram-se presentes no plenário da Comissão: Christina Aires Corrêa Lima de Siqueira Dias, Advogada Especialista da Confederação Nacional da Indústria (CNI), que convido a ocupar um local junto comigo; Luis Fernando Prado, Advogado e Professor Convidado de pós-graduação da Faculdade de Direito de Vitória.

Luis, por favor.

Obrigado, vamos deixar aqui. (*Pausa.*)

Luis, está aqui. (*Pausa.*)

Encontram-se também presentes por meio do sistema de videoconferência: Ronaldo Lemos, Diretor do Instituto de Tecnologia e Sociedade do Rio de Janeiro (ITS Rio); Matthew Reisman, Diretor de Privacidade e Política de Dados no Centro de Liderança em Política de Informação de Washington, D.C.; e Courtney Lang, Vice-Presidente de Política, Confiança, Dados e Tecnologia no Conselho da Indústria de Tecnologia da Informação de Washington, D.C., Estados Unidos.

Eu gostaria de agradecer inicialmente a presença de todos os nossos convidados, assim como todos aqueles que se encontram aqui conosco nesta sala de audiência e, também, aqueles que nos acompanham pela TV Senado e pelas redes. Lembro que esta é uma audiência pública aberta à participação de todos. As perguntas, como eu falei, podem ser enviadas pelo Portal do Senado, no e-Cidadania, ou pelo telefone, que já foi lido aqui.

É importante ressaltar que nós temos trabalhado com essa regulação do uso de inteligência artificial durante um bom tempo. Nós já fizemos mais de dez audiências públicas.

Eu sou o Vice-Presidente da Comissão, estou como Presidente em exercício aqui neste momento. Tudo começou, então, com o 2.338, proposto pelo grupo de juristas. Depois dessas mais de dez audiências, cuja maioria eu presidi, nós chegamos a um texto inicial que foi proposto na Emenda nº 1.

Quem tiver dúvida, entre no Portal do Senado, veja lá essa Emenda 1, com a ideia de se ter um texto enxuto, um texto que não atrapalhe o desenvolvimento da inteligência artificial no Brasil, um texto que não atrapalhe o uso da inteligência artificial no Brasil, mas ainda proteja, logicamente, as pessoas de forma geral do mau uso da inteligência artificial, através de uma metodologia já bastante aplicada em muitas outras áreas. Especificamente, a gente pode citar a



SENADO FEDERAL  
Secretaria-Geral da Mesa

área da aviação civil, com a Organização da Aviação Civil Internacional, outras organizações que trabalham com risco. Então, o gerenciamento de risco é feito baseado em cenário, impacto e probabilidade naquilo.

Isso aí foi colocado naquela Emenda nº 1 como forma de tabela, para que fosse essa tabela revisada periodicamente por uma comissão formada por representantes das agências reguladoras, sociedade civil, academia, setor produtivo, comunidade científica. Então, periodicamente essa tabela precisa ser revisada. Ela traz índices de 1 a 5, tanto para probabilidade quanto para impacto, e fatores contribuintes. Quem trabalha com investigação, prevenção de acidentes, sabe como a gente trabalha com fatores contribuintes distribuídos em colunas, e isso dá uma segurança operacional muito maior, assim como a segurança jurídica, pelo fato de ficar uma maneira mais objetiva da determinação de risco.

Então, esse é um dos pontos. Inclusive, imagino que nós vamos discutir aqui hoje também esse ponto. O fato é que uma legislação como essa, por se tratar de uma tecnologia que desenvolve, evolui muito rapidamente, se não tomarmos cuidado com o texto principal, corre o risco de ficar obsoleta minutos depois de ser promulgada. Então, é necessário esse cuidado, é necessário proteger as pessoas, logicamente, através dessa análise de riscos, mas também é necessário proteger o nosso mercado, ajudar o desenvolvimento das empresas aqui no Brasil, o desenvolvimento da tecnologia no Brasil, e isso, então, traz redução das restrições e também tira do texto tudo aquilo que não for concernente diretamente à inteligência artificial.

Um texto como esse... Existe uma tendência natural nossa, de forma geral, como seres humanos, com medo talvez de uma tecnologia nova, de querer introduzir dentro do texto uma série de fatores que já estão tratados, já são tratados em outras leis: o Código Civil, o Código Penal, Lei de Direitos Autorais, projeto de lei de *fake news* e assim por diante. Então, esses fatores já estão tratados lá, não convém ficarem dentro de uma legislação que trata de uma ferramenta, mesmo porque a variação ou a evolução dessa ferramenta é muito rápida. Se colocar, tender a colocar ou tentar colocar tudo dentro de uma mesma lei, ela vai ficar também bastante obsoleta rapidamente, e as implicações da inteligência artificial podem ser tratadas nas leis específicas em cada um dos setores.

A gente conta, obviamente, com as regulamentações, regulações. Então, as regulações, por exemplo, de eleições que são feitas pelo TSE tratam da inteligência artificial lá, como ela afeta as eleições lá; como ela afeta o setor de telecomunicações junto com a Anatel; como ela afeta o setor de saúde junto com a Anvisa e assim por diante, ou seja, elas tratam ali e não se pode colocar tudo isso aí dentro de uma lei, de uma ferramenta.

Então, são algumas considerações iniciais.

Novamente, eu gostaria de agradecer a todos por estarem aqui conosco.

Nós temos cinco convidados. Cada convidado terá dez minutos de exposição. E, após todos eles apresentarem, nós vamos trazer as questões feitas pelo e-Cidadania, assim como outras questões que nós tenhamos, para discussão com os nossos convidados.

E a sequência será: primeiro, remotamente, a Courtney Lang, Vice-Presidente de Política, Confiança, Dados e Tecnologia no Conselho da Indústria de Tecnologia da Informação de Washington, D.C.; depois, aqui presencialmente, Luis Fernando Prado, Advogado e Professor Convidado de pós-graduação *lato sensu* na Faculdade de Direito de Vitória; depois, de forma remota novamente, Matthew Reisman, Diretor de Privacidade e Política de Dados no Centro de Liderança em Política de Informação de Washington, D.C.; depois, remotamente também, Ronaldo Lemos, Diretor do Instituto de Tecnologia e Sociedade; e, finalmente, Christina Aires, Advogada Especialista da Confederação Nacional da Indústria.



SENADO FEDERAL  
Secretaria-Geral da Mesa

Então, vamos iniciar a nossa audiência, ouvindo os nossos convidados. Vou só checar com a mesa se está tudo certo... Está tudo certo? Está tudo conectado já lá? *(Pausa.)*

O.k. Então, remotamente, nós teremos a primeira participação: Sra. Courtney Lang, que é Vice-presidente de Política, Confiança, Dados e Tecnologia no Conselho da Indústria de Tecnologia da Informação, que fala conosco diretamente de Washington, D.C., capital dos Estados Unidos.

Está conectada já? *(Pausa.)*

**A SRA. COURTNEY LANG** *(Para expor. Por videoconferência. Tradução simultânea.)* – Muito obrigada, Senador Pontes, por organizar esta audiência e pelo convite para participar. Eu agradeço muito a oportunidade de comparecer hoje aqui perante a Comissão Temporária Interna sobre Inteligência Artificial.

Bem, meu nome é Courtney Lang. Eu sou Vice-Presidente de Política, Confiança, Dados e Tecnologia no Conselho da Indústria de Tecnologia da Informação – no Information Technology Industry Council. Eu lidero o trabalho global de política de inteligência artificial aqui no ITI.

O ITI representa empresas de todos os setores do setor de tecnologia e de todo o ecossistema de inteligência artificial, incluindo aquelas envolvidas no desenvolvimento de modelos de inteligência artificial e na implementação de aplicativos de inteligência artificial.

O ITI aplaude o esforço da Comissão para facilitar a discussão sobre políticas e regulamentação de inteligência artificial e está grato por poder contribuir com nossas opiniões e para conversa.

Apoiamos os esforços gerais do Brasil para atribuir confiança à inteligência artificial e também para promover o desenvolvimento e a implantação responsáveis da tecnologia e acompanhamos de perto o trabalho do Projeto de Lei 2.338. Vemos que o Congresso brasileiro e a Comissão Temporária têm um papel fundamental a desempenhar na proteção de consumidores e empresas, mitigando riscos previsíveis e incentivando inovações e investimentos futuros.

E, nesse sentido, eu estou muito feliz por estar aqui hoje com você, Senador, para discutir de forma mais detalhada como a Comissão Temporária pode revisar o projeto de lei para alcançar esses objetivos.

Muito embora tenhamos várias preocupações com o projeto de lei, vou me concentrar em três questões principais, se me permite, Senador: primeiramente, o ITI incentiva a Comissão a tomar medidas para revisar o projeto de lei para refletir de forma mais abrangente uma abordagem proporcional e baseada em princípios.

Vou começar aqui com a questão de direitos humanos. É uma questão fundamental para promover a confiança na tecnologia de inteligência artificial. No entanto, a abordagem atual adotada no projeto de lei que estabelece direitos individuais, no art. 5º e no art. 6º, parece que indica haver riscos significativos aos direitos humanos ao longo do ciclo de vida da inteligência artificial, o que, a meu ver, nem sempre é o caso, e ignora um contexto importante sobre o local ou sobre onde o risco geralmente se materializa durante o uso final ou a implantação dessa tecnologia.

E também, da forma que o texto da lei atualmente está estruturado, o art. 6º, que introduz os direitos sobre um sistema que produz efeitos legais relevantes ou de alto risco, e o art. 18, que impõe obrigações aos desenvolvedores de sistema de inteligência artificial de alto risco, entendemos que essa redação pode criar uma estrutura que eu caracterizaria como confusa. Por exemplo, para implementar várias das propostas baseadas em direitos na prática, vemos que as



SENADO FEDERAL  
Secretaria-Geral da Mesa

empresas podem precisar coletar mais dados pessoais, resultando em preocupações com privacidade e proteções de dados.

Além disso, eu diria que, no contexto de inteligência artificial, alguns direitos podem ser exercitáveis do ponto de vista técnico. Por exemplo, alguns desses direitos individuais que estão sendo... Eu sugeriria que seria mais útil para a Comissão remover os arts. 5º e 6º e concentrar na incorporação de obrigações no art. 18, que abordariam de forma mais proporcional os direitos individuais em cenários de alto risco, caso isso venha a ocorrer.

Se houver um desejo de buscar mais informações que possam afetar de forma mais negativa a experiência dessa pessoa, essa pessoa pode buscar mais proteção por meio dessa medida. Então, encorajá-íamos a Comissão a rever esse projeto de modo a constituir um sistema de inteligência artificial de alto risco, definir melhor este conceito: Qual o conceito de sistema de alto risco?

Bem, muito embora nós agradeçamos, apreciemos o fato de a Comissão ter tomado medidas para identificar o uso de inteligência artificial de alto risco e ter havido tentativas adicionais de detalhar mais especificamente esses casos de uso, para refletir que nem todo uso de inteligência artificial, num contexto específico, necessariamente seria de alto risco. A adição de linguagem no art. 14, por exemplo, indica que os aplicativos utilizados para os fins estabelecidos na seção 1 não são considerados de alto risco, quando não são decisivos para o resultado ou decisão, operação ou acesso a um serviço essencial. Esse texto é positivo, mas encorajamos ou sugerimos que a Comissão aprimore esse texto, de modo que o art. 14 garanta que os aplicativos que não apresentam risco não sejam capturados de forma não pretendida. Se não afeta os direitos das pessoas, isso não deveria ser classificado como de alto risco. Então, poderia se trabalhar em algum texto do art. 6º, que trata de quando o sistema de inteligência artificial é utilizado. É importante definir situações específicas em que o sistema de inteligência artificial seria classificado como de alto risco. Então seria para incluir uma linguagem que tenha como alvo específico quando o sistema é de alto risco. Então, o comitê poderia acrescentar uma linguagem que indique que, se a decisão de um sistema de inteligência artificial não impactar significativamente a segurança ou direitos humanos, além dos serviços básicos, ele não deve ser considerado como de alto risco.

O comitê deve também considerar a adição de uma linguagem de filtro em torno de casos em que o sistema de inteligência artificial não é considerado de alto risco, por exemplo, a proposta do Artigo 6 da Lei da União Europeia sobre Inteligência Artificial. O sistema de inteligência artificial se destina a executar uma tarefa processual restrita.

Enfim, esse tipo de linguagem reflete até certo ponto aquela introduzida na lei de inteligência artificial que vocês estão propondo. Nós apoiamos essa linguagem, porque ela acrescenta ainda mais especificamente aos casos em que o sistema de inteligência artificial não deve ser considerado de alto risco. No entanto, nós observamos que a Comissão Europeia ainda deverá emitir diretrizes adicionais até fevereiro de 2026, incluindo exemplos práticos de sistemas de inteligência artificial de alto risco e de baixo risco. E o Brasil também pode precisar criar um processo pelo qual diretrizes adicionais possam ser emitidas para, por exemplo, definir, criar um processo para envolver as partes interessadas de modo a refinar ainda mais o art. 15 – isso entra como uma sugestão –, de modo a avaliar se um sistema de inteligência artificial é ou não é de alto risco.

E o critério atual... Essa habilidade atualmente está, de certa forma, ambígua, para que as empresas possam compreender como cumprir essa previsão de se esse sistema que eles estão produzindo é de alto risco ou não é.





SENADO FEDERAL  
Secretaria-Geral da Mesa

Eu já entendi que tem várias conversas sobre qual critério que deveria ser estabelecido e tem um desejo de estabelecer um critério no art. 15 desse projeto de lei, envolvendo todos os atores envolvidos, incluindo participantes da indústria e também as partes que opinam no critério, por exemplo... Pode tomar a fórmula de uma audiência pública, com seis meses de prazo, para que as partes possam apresentar suas contribuições ou apresentar para cumprir os seus requisitos, se o desejo for manter esse critério de avaliação de risco do art. 15 atualmente apresentado.

E, por fim, nós sugerimos que o comitê...

De modo geral, por fim, nós encorajamos e sugerimos que o comitê possa definir as obrigações em torno dos sistemas de inteligência artificial de modo geral. Deve-se considerar a aplicação de medidas de governança aos modelos de GPAl, modelos mais avançados, em vez de sistemas de inteligência artificial de uso geral, que, se baseados na definição atual, abrangem uma ampla variedade de sistemas que podem ser utilizados para qualquer número de protótipos, incluindo aplicações de baixo risco.

Reconhecemos que modelos de inteligência artificial de uso geral podem apresentar riscos exclusivos e concordamos que os provedores desses modelos podem tomar medidas para aumentar a segurança, a transparência e reduzir os riscos. Mas é importante reconhecer que a gestão eficaz de riscos é, em última análise, específica do contexto e, em muitos casos, depende de circunstâncias da implantação. E certamente o objetivo é ampliar a implantação desses modelos, mas reconhecemos que esses modelos... Ao desenvolver medidas de governança, o comitê deve reconhecer que nem sempre é possível para um provedor prever cada uso individual desse modelo, e os requisitos devem ser calibrados para o nível de risco do modelo, ser praticáveis e se estender apenas para o que os provedores possam razoavelmente abordar durante o projeto e o desenvolvimento.

E nós temos outras sugestões específicas, que eu poderei compartilhar com vocês de formas específicas. Eu tenho preparado um artigo a respeito de tudo isso, eu poderei compartilhar com o senhor, Senador, de forma mais específica, as sugestões específicas. A minha equipe também estará no Brasil na semana que vem, e, se tiver mais oportunidades em que possamos nos reunir e contribuir, estarei à disposição.

Novamente eu agradeço pela oportunidade de estar aqui. Eu estarei à disposição para responder a qualquer questão que vocês poderão apresentar.

Muito obrigada.

**O SR. PRESIDENTE** (Chico Rodrigues. Bloco Parlamentar da Resistência Democrática/PSB - RR) – Quero agradecer a participação da Courtney Lang, que obviamente tinha iniciado a sua apresentação ainda na Presidência do Senador Marcos Pontes e que obviamente, ao longo das apresentações, deverá receber realmente perguntas de algum Senador ou de algum participante desta importante reunião.

Continuando a lista de convidados da Comissão, eu quero chamar Luis Fernando Prado, Advogado e Professor convidado de pós-graduação *lato sensu* na Faculdade de Direito de Vitória.

V. Sa. dispõe de dez minutos.

**O SR. LUIS FERNANDO PRADO** (Para expor.) – Boa tarde a todas e a todos.

Eu queria cumprimentar o Senador Chico Rodrigues pela Presidência aqui da nossa sessão. Eu queria agradecer também o convite feito pelo Senador Astronauta Marcos Pontes e parabenizar pela iniciativa. Tem sido claro e evidente que, quanto mais tempo a gente tem de debater o texto do PL 2.338, mais as novas versões que surgem de texto acabam refletindo



SENADO FEDERAL  
Secretaria-Geral da Mesa

melhorias e um aumento da maturidade também da redação. Nesse sentido, eu queria parabenizar e cumprimentar também o Senador Eduardo Gomes pelo trabalho, o árduo trabalho feito na relatoria dessa matéria, uma matéria supercomplexa, superdisruptiva – para usar a palavra da moda –; e também cumprimentar os meus colegas da Comissão de Juristas, que também fizeram um trabalho incansável aqui de propor a redação inicial que a gente hoje discute.

Dito isso, Senador, o meu papel aqui, minha missão, é trazer alguns pontos em que, ao meu ver, ainda cabem melhorias em relação ao PL 2.338. São pontos que fazem, na minha opinião, inclusive nosso PL ser mais rigoroso do que a regulamentação europeia sobre o tema, que é considerada uma das regulamentações mais rigorosas mundo afora que a gente tem e que serve de base aqui para a nossa proposta legislativa; mas em alguns pontos a gente talvez tenha importado obrigações, e não as exceções trazidas ali pela regulamentação europeia.

E focando então nos pontos que eu gostaria de trazer aqui, o primeiro deles é relativo a direitos autorais. Eu acho que os arts. 60 e seguintes da proposta atual de texto contam com espaço de melhoria ainda na sua redação para torná-la mais equilibrada. E aqui eu faço algumas perguntas. O que a gente quer enquanto sociedade quando a gente olha para a IA? A gente quer sistemas mais ou menos seguros? A gente quer sistemas bem ou mal treinados? A gente quer sistemas que alucinam ou que não alucinam? A resposta ou as respostas para as perguntas passam por a ampla possibilidade de a gente treinar os sistemas com base em dados publicamente disponíveis, dados legalmente disponíveis. Quando a gente fala de treinamento de IA, a gente não está falando só de lucro, a gente está falando de segurança e confiabilidade dos sistemas também por parte das organizações, por parte da sociedade, o que deveria ser o principal norte aqui da nossa regulamentação sobre o tema.

E aqui eu faço um paralelo com os navegadores. Na minha opinião, se a gente tivesse uma regra tão rígida em relação a direitos autorais quando os navegadores – e mais do que os navegadores, os buscadores aqui – e as ferramentas de busca da internet se popularizaram, talvez a gente nem sequer teria os motores de busca como a gente conhece hoje, que são tão úteis à sociedade e tão essenciais.

Então, eu acho que aqui cabe mais espaço para a gente achar uma redação equilibrada em relação ao tema que possibilite o uso de obras legalmente disponíveis, pelo menos quanto à análise computacional dessas obras – claro, respeitadas as condições que também preservem os direitos de titulares, de direitos do autor.

Então, o meu primeiro ponto de contribuição é especificamente sobre esse eixo aqui dos direitos autorais.

Avançando um pouco mais aqui e ainda falando em pontos que a gente poderia importar também da experiência europeia para tornar a nossa regulação mais equilibrada, tem um ponto bastante importante que é a extensão da aplicação da nossa legislação. A nossa legislação, de acordo com a proposta que hoje está prevista, ela se aplica inclusive à fase de teste, onde as falhas e os erros são esperados e encorajados – e que bom que eles aparecem nessa fase de teste antes da colocação em mercado. É claro que, se eu estou desenvolvendo um sistema de IA, eu tenho que desde já observar os ditames regulatórios para, quando ele entrar em mercado, ele ser um sistema em conformidade com a regulação. Mas me parece muito duro e muito um contrassenso, até, que a empresa ou uma organização que esteja desenvolvendo esse sistema de IA possa ser punida e responsabilizada por erros e falhas que aconteçam na fase de teste. Então me parece que a regulação deveria abrir margem para que a gente comemore erros e falhas que apareçam na fase de teste antes de a gente colocar o sistema em mercado.



SENADO FEDERAL  
Secretaria-Geral da Mesa

E aqui eu deixo como inspiração talvez o texto do Artigo 2º, o item 8, do AI Act, que deixa claro que a regulamentação não se aplica, pelo menos na sua parte de punição e responsabilização, aos sistemas em fase de teste, ambiente de teste que ainda não foram colocados em mercado.

Outra exceção importante trazida pela legislação europeia e não refletida ainda na proposta brasileira diz respeito aos sistemas disponibilizados sob licenças de código aberto. Então, fica minha contribuição também para a gente refletir se não seria o caso de a gente incorporar essa exceção também no nosso texto do PL 2.338.

Avançando um pouco mais, eu queria destacar a obrigação de explicabilidade que aparece no PL 2.338. A gente tem dois tipos de obrigação no PL 2.338 quando a gente está falando de transparência: requisitos objetivos e claros de transparência, e também a explicabilidade, que aparece três vezes no texto legal. Explicabilidade e transparência me parecem... Na verdade, a explicabilidade na intenção original do texto seria algo mais do que a transparência, que já está prevista. O problema é que a explicabilidade é um conceito jurídico aberto, amplo, indeterminado e passível de interpretações bastante subjetivas. Então me parece que aqui faria mais sentido a gente focar em requisitos claros e objetivos de transparência, que a gente vá conseguir fiscalizar, cobrar e cumprir esses requisitos, do que usar termos bastante amplos que vão deixar a legislação sujeita às mais diversas interpretações e que vão desafiar o estado da arte inclusive, porque os sistemas de IA caminham para serem cada vez mais complexos e talvez menos explicáveis – essa é a minha experiência, pelo menos no campo digital. No campo de direito digital, toda vez que eu converso com os times de tecnologia, eles me comentam: "Luis, se eu soubesse exatamente por que o sistema está fazendo isso, se eu pudesse explicar tudo o que o sistema está fazendo, eu teria desenvolvido um código e não estaria usando IA aqui para desenvolver esse sistema".

Então me parece, aqui, que a gente deveria focar em requisitos claros e objetivos de transparência e não em termos ambíguos e bastante abrangentes, como a explicabilidade, que nem sequer está definida expressamente no texto, na proposta que a gente tem na mesa.

Avançando ainda um pouco mais, aqui já finalizando minhas contribuições iniciais, eu queria tratar das avaliações que o nosso projeto de lei traz.

A gente tem basicamente dois tipos de avaliações previstas na proposta: uma relativa às avaliações preliminares, que são aquelas avaliações que precisam ser feitas em todos os casos; e a outra, a avaliação de impacto algorítmico.

Em relação às avaliações preliminares, elas são exigidas de todo e qualquer sistema de IA e de todos os agentes, ou seja, tanto o desenvolvedor quanto o aplicador do sistema. E não é de qualquer tipo de avaliação que a gente está falando aqui: é uma avaliação formal, que precisa ser registrada, documentada e armazenada por cinco anos por todos os agentes. Qual é o meu ponto aqui? Sistemas de IA amanhã vão ser sinônimo de *software*. Qualquer ferramenta vai ter um sistema de IA embarcado ou vai ser um próprio sistema de IA, e, se a gente precisar fazer, formalizar, documentar uma avaliação de cada *software* que a gente usa no nosso dia a dia empresarial, a gente vai ter que parar as empresas, colocar um setor específico para isso, atraindo um custo regulatório gigantesco e um ônus, obviamente, de armazenar essa documentação por cinco anos, manter isso atualizado. Então me parece uma obrigação talvez bastante excessiva para a gente exigir – é uma obrigação que não aparece expressamente inclusive na legislação europeia –, me parece que esse é um ponto bastante excessivo, em que cabe algum tipo de flexibilização aqui para a gente conversar um pouco mais sobre a realidade prática das empresas.





SENADO FEDERAL  
Secretaria-Geral da Mesa

E aqui eu faço um paralelo com a legislação de proteção de dados, porque alguém pode estar se perguntando: "Como eu vou saber o risco do sistema se eu não faço uma avaliação prévia, formal, documentada?". Bom, esse desafio a gente já tem hoje na legislação de proteção de dados. A gente tem obrigações na LGPD que se aplicam para tratamento de dados de alto risco, e não necessariamente o legislador da LGPD previu que eu precisava fazer uma avaliação formal, armazenar por cinco anos uma avaliação prévia de cada tratamento de dados que eu faço. Então, eu acho que cabe aqui uma flexibilidade para que as organizações decidam qual é o melhor formato dessa avaliação, como fazer essa avaliação de risco, e não necessariamente a lei colocar essa obrigação tão abrangente e tão rígida a ponto de a gente ter que registrar e documentar isso por cinco anos.

Por fim, Senador, tem um ponto também que me chama a atenção em relação a esse assunto das avaliações, e aqui eu passo para a avaliação de impacto algorítmico. Qual é o problema aqui, a meu ver? A gente tem um dispositivo hoje na proposta de lei que a gente está comentando que diz que as conclusões das avaliações de impacto algorítmico devem ser publicadas. E eu fico bastante preocupado com esse dispositivo, porque uma avaliação de impacto algorítmico vai evidenciar talvez fragilidades, que precisam obviamente ser corrigidas e endereçadas, mas vai evidenciar talvez riscos aqui. E, ao expor essas conclusões, eu não tenho segurança de que isso é bom para a sociedade, de que a gente vai ter ganhos relevantes com isso. Talvez a gente esteja na verdade expondo ainda mais os nossos sistemas, talvez a gente esteja contribuindo para que essas informações sejam utilizadas por agentes maliciosos em relação ao funcionamento dos nossos sistemas. Então, a minha conclusão em relação a esse ponto é que o bônus dessa transparência em relação à conclusão de uma avaliação de impacto algorítmico não compensa o ônus nesse caso.

Eu encerro minha fala, deixo aqui os meus minutinhos excedentes, devolvo a palavra para o Senador, agradecendo-lhe e parabenizando, mais uma vez, o Senado Federal por este espaço de debate tão essencial e que certamente vai fazer a nossa redação do PL 2.338 ser evoluída e ganhar em maturidade.

Muito obrigado, Senador.

**O SR. PRESIDENTE** (Chico Rodrigues. Bloco Parlamentar da Resistência Democrática/PSB - RR) – Como todos que estão acompanhando esta sessão estão vendo, essa instrução que está sendo dada a esse Projeto de Lei 2.338, de autoria do Presidente da Casa, Rodrigo Pacheco, é extremamente complexa, é extremamente atual e, nas suas transversalidades, demonstra a necessidade do aprofundamento desses estudos, para que o projeto saia pronto e acabado e se mantenha em completa reciclagem, porque, quando se fala de IA, não tem limite, não tem fim. Todos sabem que a cada nova descoberta, a cada novo sistema, você tem que sair fazendo adequações, porque o que sai da imaginação humana, na verdade, não tem limite. O céu é o limite e ninguém sabe onde é o limite do céu.

Então é muito complexo, é muito importante aqueles que se dedicam, que se debruçam sobre essas questões.... A gente acompanhava anteriormente a apresentação da Courtney Lang, mostrando, inclusive, um organismo importantíssimo, o ITI. E mesmo assim, lá nos Estados Unidos, o que eles falaram para nós é que eles estão numa fase vegetativa de descoberta, todo dia tem um novo modelo, enfim.

Então, continuando essa lista de convidados, passamos a palavra para o Matthew Reisman, Diretor de Privacidade e Política de Dados no Centro de Liderança em Política de Informação de Washington.

O seu tempo é de dez minutos. *(Pausa.)*



SENADO FEDERAL  
Secretaria-Geral da Mesa

Está conectando?

**O SR. MATTHEW REISMAN** (Para expor. *Por videoconferência. Tradução simultânea.*) – Presidentes em exercício Senador Pontes e Senador Rodrigues, boa tarde.

Meu nome é Matthew Reisman e sou Diretor de Privacidade e Política de Dados no Centro de Liderança em Política de Informação. Gostaria de agradecer ao Senador Astronauta Pontes e ao comitê por me convidarem para testemunhar hoje. Obrigado também por me permitirem testemunhar em inglês. Vou continuar nessa língua.

CIPL é uma organização com escritório em Washington e em Londres, que trabalha com o uso benéfico de dados. A inteligência artificial é o foco da CIPL, inteligência artificial e proteção de dados, desde que nós temos vários parceiros que se dedicam à proteção de dados e inteligência artificial no mundo inteiro, inclusive no Brasil.

Gostaria de compartilhar com vocês algumas conclusões recentes sobre pesquisas e também oferecer algumas recomendações específicas para melhoria no projeto de lei que está sendo discutido. No começo deste ano, a CIPL publicou um documento documentando as melhores práticas entre 20 projetos e classificou esses projetos em vários elementos a respeito dentro de um *framework* ou de um marco de avaliação da CIPL. Analisa avaliação de risco, procedimentos, transparência, treinamento, alerta, resposta e execução ou abordagem. Esse é um modelo para desenvolver uma abordagem abrangente de resposta e de transparência para o trabalho com inteligência artificial. Isso é consistente com esse marco de avaliação. Muitas organizações já enxergam a governança de IA como importante para a sua avaliação. Esses programas exigem tempo e recursos, e muitas organizações estão se dedicando a isso.

Em termos de Smart IA, na regulação de inteligência artificial foram publicadas dez recomendações para a regulação de inteligência artificial. Esses relatórios tratam de três... A parte principal, incluindo a tecnologia e o quadro de avaliação, e o poder desses indivíduos dentro desse marco de avaliação. No próximo nível, a parte demonstrável e de responsabilização das práticas de inteligência artificial, como fator de mitigação desses fatores de execução ou de supervisão ou fiscalização. Por fim, eu trataria dessa parte de avaliação de inteligência artificial.

Na sua avaliação ou redação mais recente, nós vemos algumas áreas para melhoria, e, como já foi dito aqui pelos participantes anteriores, na lei de inteligência artificial na União Europeia, que trabalhou no desenvolvimento dessa legislação, o que foi produzido foi algo muito complexo, e também chegamos a uma conclusão de que precisamos de tempo para estabelecer um requisito legal, de forma prática, para que seja aprovado.

E os participantes também precisam lidar com as tensões, com as legislações existentes, e precisamos pensar também no valor de ter uma solução mais prática e evitar conflitos com as legislações existentes. E, de forma específica sobre o projeto de lei aqui em questão, precisamos tratar das questões que são excessivas e também sobre alto risco. Essas questões deveriam ser flexíveis, como resposta a essas questões envolvendo a apresentação. E essas questões de alto risco deveriam ser rebatidas, como a Courtney trouxe no seu discurso.

Essa parte dos códigos de boas práticas também precisa estar muito bem definida dentro do projeto, e esses códigos precisam ser tratados como atos de boa vontade e de fiscalização.

O conselho permanente de inteligência artificial proposto trabalharia ao longo de diferentes agências e faria todos os esforços para que essa instituição pudesse trabalhar, e trabalhar em coordenação com todos os atores. E, ao trabalhar com esses *sandboxes*, eles trabalham de forma a proteger as pessoas enquanto estão desenvolvendo inovação, mas a lei não apresenta grandes incentivos para o uso desses *sandboxes* ou ambientes temporários de teste. E é importante desenvolver algumas considerações sobre esses processos de fiscalização ou de



SENADO FEDERAL  
Secretaria-Geral da Mesa

supervisão e o impacto de avaliação ou de impacto. E a redação atual é que a avaliação do algoritmo por parte dos desenvolvedores e dos empregadores ou dos usuários dentro das diferentes instituições que são utilizadas... E precisamos de uma base de dados pública para a inteligência artificial e, se essa base de dados for criada, precisamos produzir acesso com as devidas proteções para preservar a confidencialidade dos dados e para evitar a exploração de diferentes agentes. O uso desses recursos é um desafio, porque vai exigir grandes considerações sobre as questões aqui.

E sabemos que é importante garantir que os reguladores que implementam a regulação também precisam ser educados sobre inteligência artificial. Quando eles entenderem a tecnologia e os recursos, eles vão estar mais capacitados para enfrentar os possíveis riscos e desafios. E a CIPL acredita que esse recurso com requisitos robustos e com programas de cumprimento da legislação robustos e claros...

Com isso, eu concluo aqui a minha apresentação e me coloco à disposição para qualquer esclarecimento durante a fase de perguntas.

**O SR. PRESIDENTE** (Chico Rodrigues. Bloco Parlamentar da Resistência Democrática/PSB - RR) – Quero agradecer ao Matthew Reisman pela sua importante colaboração nesta Comissão Temporária Interna sobre Inteligência Artificial.

Continuando a lista de convidados, ainda de forma remota, eu convido o Ronaldo Lemos, Diretor do Instituto de Tecnologia e Sociedades (ITS) do Rio.

Você dispõe de dez minutos.

**O SR. RONALDO LEMOS** (Para expor. *Por videoconferência.*) – Muito obrigado, Senador.

Peço só um minuto, porque neste momento o interfone tocou aqui. Se puderem aguardar um minuto, só para o barulho parar. Um minuto.

(Pausa.)

Peço desculpas pela interrupção.

Bom, primeiramente, gostaria de saudar o Senado Federal, na figura do Senador Marcos Pontes, que gentilmente me fez o convite para estar aqui hoje. É um prazer imenso poder participar desse debate sobre um tema tão importante, que é a construção da lei de inteligência artificial do Brasil.

Aproveito para me apresentar. Meu nome é Ronaldo Lemos, eu sou advogado, trabalho com questões envolvendo políticas tecnológicas há mais de 20 anos – o tempo voa! – e participei da construção do marco civil da internet desde o início, quando eu formulei a ideia sobre o projeto de lei e também na articulação do marco civil, na gestão da consulta pública que fez com que o marco civil fosse construído e, depois, na aprovação, intensa articulação e debate no Congresso Nacional até a aprovação final, em 2014, dessa lei tão importante, que é a base da regulação tecnológica aqui no Brasil.

Dito isso, eu vou fazer minha apresentação dividida em duas partes: na primeira parte, eu vou apresentar alguns problemas que eu vejo na questão do projeto de lei de inteligência artificial, tanto de processo quanto de conteúdo; e, na segunda parte, eu vou apresentar algumas sugestões que eu gostaria de submeter à apreciação do Senado Federal com relação a como é que esse projeto pode ser aperfeiçoado.

Então, eu vou começar primeiro falando e apontando aqui os problemas.

A primeira coisa que eu acredito que vale a pena notar é que, como a minha experiência no processo legislativo abrange a experiência do marco civil, esse é um processo que está tendo um grau de participação significativamente menor do que, por exemplo, o processo que levou à construção do marco civil. Eu acho que essa audiência pública aqui é excelente, todas as que



SENADO FEDERAL  
Secretaria-Geral da Mesa

foram realizadas foram excelentes também, os pedidos de contribuição, mas, como alguém que circula pelo ambiente de tecnologia no Brasil, a maioria das pessoas no nosso país hoje ainda não sabe e nem teve a oportunidade de participar da formulação da nossa lei de inteligência artificial. Inclusive, pessoas que trabalham com tecnologia ainda não sabem que existe até um debate em curso e que a lei está em estágio tão avançado para ser aprovada. Então, primeiro ponto, eu reforçaria o aspecto da participação e abriria mais oportunidades, inclusive setoriais, para a participação ampla neste projeto.

O segundo ponto que eu acho que vale enfatizar é a questão da transparência. O nosso processo legislativo aqui teve algumas falhas, por exemplo as contribuições que foram oferecidas por meio de consulta pública não foram publicadas pelo Senado Federal, foi preciso entrar com pedido de LAI (Lei de Acesso à Informação), para que o Senado publicasse as contribuições que foram recebidas por consulta pública, muito embora a LAI (*Falha no áudio.*) ... sendo publicadas. Então essa transparência é interessante e importante para a sociedade saber até o que está sendo contribuído e quais as preocupações com relação ao projeto.

Terceiro problema que eu gostaria de apontar é com relação à matriz do nosso projeto de lei. O nosso projeto está, na prática, copiando a matriz do modelo europeu de inteligência artificial. O modelo europeu, eu acho que é um modelo importante, bem feito, mas eu acho que o Brasil não deveria ter o compromisso de copiar esse modelo. Eu acho que, quando fizemos o marco civil da internet, ele não foi copiado de nenhum lugar; ele é um projeto de lei que foi construído de acordo com as necessidades e contribuições aqui do Brasil. Então essa cópia do modelo europeu não me parece ser o caminho totalmente desejável. Inclusive, veio ao Brasil há alguns meses um dos autores da lei europeia, o Gabriele Mazzini, que deu entrevista para jornais, como *Folha de S.Paulo* e outros, e a mensagem dele era muito clara, ele dizia: "Não copiem o modelo europeu". O modelo europeu é basicamente construído com as características da União Europeia, é uma região que tem um PIB muito maior do que o PIB brasileiro, são vários países, tem uma dinâmica completamente diferente. Então, esse ponto me parece importante.

E o outro problema da contribuição do modelo europeu é que, ao transplantarmos essa matriz para o Brasil, a gente esqueceu de copiar a parte propositiva do modelo europeu, que é a parte que tem a sua integração no fomento aos setores produtivos. Então, o modelo europeu não é só para tratar da questão do risco, ele também propõe uma série de medidas para que a Europa se torne altamente competitiva, inclusive com designação de investimentos e agências para acompanhamento disso tudo, e essa parte acabou não ingressando na nossa lei.

Na minha visão, a lei brasileira precisaria ter um tripé para que ela fosse uma lei que atendesse aos temas gerais de interesse do país. O primeiro elemento desse tripé é a questão do risco, e essa nós já temos. A lei toda basicamente é dedicada a tratar da questão do risco da inteligência artificial. E o que está faltando? Está faltando outro eixo que tratasse da questão do trabalho: como é que a gente vai capacitar trabalhadores no Brasil, como é que a gente vai aumentar a nossa produtividade com relação à inteligência artificial, como é que a gente vai se preparar os brasileiros e as brasileiras para o trabalho do futuro? Isso não consta no projeto de lei como poderia constar. E, por fim, o terceiro elemento desse tripé é a questão da concorrência e da competitividade. A concorrência é para evitar concentração econômica que a inteligência artificial pode trazer de forma, inclusive, nunca vista. Então, esse ponto é importante de ser abrangido. E competitividade é fomentar a indústria brasileira, o agro brasileiro e outros setores para que sejamos competitivos, inclusive internacionalmente, em inteligência artificial.

Para concluir a parte dos problemas, se o projeto de lei fosse um time de futebol, é como se a gente estivesse escalando só zagueiros para jogar nesse time. E, como a gente sabe muito



SENADO FEDERAL  
Secretaria-Geral da Mesa

bem, um time de futebol que só tem zagueiro não ganha. A gente precisa ter zagueiros, precisa ter meio de campo e precisa ter atacantes. Essas outros dois pilares não estão presentes no projeto hoje, porque essencialmente ele só trata de um tema, que é o tema do risco.

Eu encerro aqui a parte dos problemas e passo para a parte das sugestões.

A primeira sugestão que eu gostaria de tratar, já dentro dessa ideia de a gente ter propostas e elementos propositivos dentro do projeto, é a questão da energia. Inteligência artificial vai consumir energia de forma abundante. A estimativa é que 4,5% da energia do planeta vão ser consumidos por *data centers* nos próximos anos. O Brasil é o único país que tem uma matriz energética com 93,1% de energia limpa. O segundo lugar é muito distante do Brasil. O que está acontecendo agora é que as empresas de tecnologia estão aumentando as suas emissões de carbono, e todas estão buscando fontes de energia renovável. Na minha visão, o Brasil pode enriquecer se ele souber vender certificados de energia renovável globalmente e também apostar em projetos como o do hidrogênio verde, no que até o Senado teve um papel essencial, pois acabou de aprovar o marco do hidrogênio verde.

Então, o Brasil tem a possibilidade de enriquecer vendendo energia para os *data centers* de inteligência artificial. Essa é uma oportunidade histórica para o país, e o projeto de lei poderia tratar disso. Isso é um elemento de desenvolvimento nacional que está faltando quando a gente pensa em inteligência artificial do ponto de vista legal e regulatório.

A segunda sugestão que eu vou fazer é que a gente precisa fomentar consórcios internacionais de aquisição de insumos tecnológicos. O Brasil hoje está sozinho no desenvolvimento de inteligência artificial e, tal como já fizemos, por exemplo, com vacinas – o Brasil participa de vários desses consórcios para aquisição de vacinas –, é muito importante que o Brasil desenvolva e participe usando o seu poder de *procurement*, o seu poder de compras públicas associado ao de outros países para aquisição de insumos que permitam inclusive treinar e desenvolver modelos de inteligência artificial. Isso também não está previsto no projeto, e eu entendo como sendo um elemento essencial para o futuro do país na sua autonomia, soberania e possibilidade de competição tecnológica.

A terceira sugestão é a integração entre o Plano Brasileiro de Inteligência Artificial com o projeto de lei. O projeto de lei é muito diferente do PBIA. O PBIA é um plano que realmente tem olhado para frente, olhado para questões como aplicação da inteligência artificial para a área de saúde, serviços públicos, entre outras, mas, quando se analisam o projeto de lei e o plano, parece que eles foram construídos em países diferentes. Por exemplo, o plano brasileiro fomenta a inteligência artificial na saúde e a lei diz que saúde é uma atividade de alto risco, que merece um escrutínio e um peso regulatório absolutamente excessivo. Então, de um lado, a gente aponta para uma questão fundamental que é a saúde e como ela se integra na inteligência artificial e, de outro, a gente segura o desenvolvimento de aplicações nessa área, na sua máxima segurança, fazendo com que toda aplicação de saúde, ou boa parte dela, seja de alto risco.

Então, eu acho que existem contradições entre o Plano Brasileiro de Inteligência Artificial e o projeto de lei. Eu gostaria que a gente tivesse feito o contrário, que a gente tivesse feito o Plano Brasileiro de Inteligência Artificial e desenvolvido o projeto de lei a partir dele, porque isso nos teria dado um norte de aonde queremos chegar. E eu acho que esse é o ponto fundamental para a gente fazer qualquer política tecnológica. Nenhum vento ajuda quem não sabe a que porto veleja. Então, resolver essas contradições me parece importante.

Outro ponto é o desenvolvimento de capacidades locais. A gente precisa fazer com que a inteligência artificial seja aplicada na indústria, no agro e em diversas outras capacidades setoriais aqui do Brasil





SENADO FEDERAL  
Secretaria-Geral da Mesa

Para terminar, o que eu gostaria era a criação de uma lei de inteligência artificial que fosse original, brasileira e à altura das capacidades do país. Para isso a receita é: participação, escuta de toda a sociedade e mobilização. Quanto mais gente participa e contribui para o plano – como esta audiência aqui demonstra – mais ideias a gente vai coletar, mais possibilidades de atuação a gente vai ter, e vamos ter uma lei de IA comprometida com os interesses nacionais e com as capacidades que o Brasil também tem de desenvolver e competir nessa área.

Muito obrigado, Senadores. Obrigado aos colegas.

Eu encerro aqui a minha participação.

**O SR. PRESIDENTE** (Chico Rodrigues. Bloco Parlamentar da Resistência Democrática/PSB - RR) – Quero agradecer ao Ronaldo Lemos, Diretor do Instituto de Tecnologia e Sociedade – ITS, pela sua rica participação, obviamente estimulando o debate com suas ideias e suas sugestões.

A próxima convidada é a Sra. Christina Aires Correa Lima de Siqueria Dias, Advogada Especialista da Confederação Nacional da Indústria (CNI).

**A SRA. CHRISTINA AIRES CORREA LIMA DE SIQUERIA DIAS** (Para expor.) – Obrigada, Senador Chico Rodrigues. Também agradeço ao Senador Astronauta Marcos Pontes, em nome da CNI, a oportunidade de abrir esse debate. Realmente, a gente entende que é um projeto muito complexo, que envolve, dentro da indústria, cadeias de valores de todos os tamanhos, ou seja, de toda complexidade, e que nos deu a oportunidade, inclusive, de agora fazer uma rodada muito técnica, trazendo todos os setores. Alguns setores públicos estão aderindo, querendo ser ouvidos, universidades, que desenvolvem IA no Brasil., e por quê?

Aqui eu já começo dizendo que a gente está de acordo com tudo o que foi dito. Eu fiquei muito feliz de ver que os pontos levantados por todos os debatedores que me precederam são exatamente os pontos de preocupação da CNI que a gente entende que precisam de melhoria no projeto.

Então, vou começar aqui trazendo alguns pontos de preocupação e alguns caminhos que nós estamos levando aos setores, para ver se realmente são os mais adequados para uma regulação de IA que obviamente seja responsável, respeite todos os princípios, mas que não nos tire da rota de desenvolvimento nem do mercado global desse sistema

Alguém colocou, acho que foi você, que eles vão virar *software*. Eles já são *softwares* banais na indústria. A indústria já opera com sistema de IA para competitividade e produtividade, e, sem esse sistema, a gente pode até desindustrializar o país, porque nós não vamos ter capacidade de competir com o mundo afora.

Então, nós temos sistemas que não têm risco nenhum, que não interagem com o ser humano, que não têm dados pessoais, que são máquina com máquina. E, de repente, a preocupação de toda a indústria é de se ver com mais obrigações de governança e etc. que não têm nenhum risco a direitos fundamentais ou ao ser humano. Aliás, pelo contrário, ela vai trazer problemas para o ser humano. Tem empregabilidade, insegurança e insegurança inclusive de produtos industriais que são utilizados para a segurança do produto que vai chegar na mão do consumidor final.

Então, é com essa preocupação muito grande e muito séria que a CNI vê com muitos bons olhos essa parada da Comissão para ouvir os setores e ouvir a sociedade.

Aí já vamos entrando mais no texto: uma das grandes preocupações da CNI é o desenvolvimento científico e tecnológico. Nós já colocamos isso em todas as manifestações, já colocamos para os Senadores. Então, uma lei que começa dizendo que vai regular a concepção



SENADO FEDERAL  
Secretaria-Geral da Mesa

e o desenvolvimento da tecnologia nos preocupa muito. E aí acho muito interessante, porque todos colocaram – e a gente também colocou isso...

Uma coisa que às vezes está sendo confundida no debate não é a regulação do desenvolvimento; é o momento da exigência dessa regulação.

E a Europa até deixa isso muito claro, porque ela fala: não devem ser regulados a concepção, o desenvolvimento, as fases de desenvolvimento; a regulação vai incidir a partir da entrada do produto no mercado. Do contrário, acontece o que falava nosso Diretor de Inovação: eu vou me sentar para programar com um regulador do lado. Ninguém faz isso, ninguém no mundo, nem na Europa. Então, essa falta de explicitação desse tipo de coisa no projeto nos remete a uma regulamentação que as pessoas falam que está mais restritiva até mesmo do que a da União Europeia.

Então, nós entendemos que nas exceções precisam vir realmente essas questões de eximir o desenvolvimento científico e tecnológico não só para os sistemas exclusivos para o desenvolvimento científico e tecnológico, mas para o desenvolvimento de todos os sistemas antes de entrar no mercado.

Outra questão importantíssima, que aí também foi colocada agora, é a questão do fomento. A CNI vem colocando isso. E nós desconsideramos toda uma legislação que nós já temos – a Lei de Inovação, os arts. 218 e 219 da Constituição – que coloca o mercado de inovação como um patrimônio nacional, porque é exatamente com a inovação que nós vamos ter desenvolvimento social e econômico neste país. Não adianta achar que o desenvolvimento social e econômico vai ficar fora do desenvolvimento científico e tecnológico. É exatamente o oposto. Então, a luta da CNI é para que nós continuemos competitivos, gerando emprego, renda e impostos, até, para o Governo, para a gente poder ser um país que consegue arcar com toda a nossa carga social, que também está lá no art. 6º. Então, um projeto que pretende dar direitos fundamentais e direitos sociais que dependem de dinheiro para prestações positivas do Estado e depois mata a inovação e a tecnologia, que é exatamente como o Estado vai conseguir se desenvolver e gerar renda e impostos, não casa com nenhum dos nossos princípios constitucionais, muito menos com o 218 e o 219.

E aí a Lei de Inovação já traz uma série de medidas para fomentar a inovação que nós achamos que têm que entrar no texto. E aí, como medidas de fomento à inovação, nós vemos até uma flexibilidade regulatória – alguém colocou isso aqui também; já não lembro quem, porque eu estava concordando com todo mundo, e já não lembro quem foi. Mas, enfim, a regulação pode funcionar como uma barreira regulatória – e muitas delas a gente vai até a OMC para reclamar dos outros países quando nos colocam – ou ela pode funcionar como um estímulo ao desenvolvimento de alguma atividade econômica. Então, nós entendemos que aqueles atores e aquele sistema nacional de inovação que é estimulado e fomentado na Lei de Inovação devem ter uma flexibilidade regulatória nessa lei.

E, como colocado pelo Senador Marcos Pontes, essa lei deve ser uma ferramenta. Quando você lê sobre inovação e como regular inovação e novas tecnologias, o que se coloca? Quais são as ferramentas para regular novas tecnologias? E aí nós vemos *sandbox* regulatório, você vê autorregulação regulada, uma série de ferramentas. A lei também é uma ferramenta, mas a lei não deve ser uma ferramenta para tudo e de forma geral, porque, como também foi colocado, ela não vai deixar que a gente acompanhe a normal evolução do mercado, ainda mais numa tecnologia dessa em que às vezes as atualizações de uso são semanais. O que a CNI entende mais adequado? Você...



SENADO FEDERAL  
Secretaria-Geral da Mesa

E outra coisa: esses desenvolvimentos da Lei de Inovação são para aquelas tecnologias desenvolvidas no país, nos centros de inovação, para solucionar problemas brasileiros, para solucionar questões de produtividade nacional.

Nós temos vários projetos em universidades públicas, que são fantásticos, que trazem a indústria com problemas para os estudantes desenvolverem. Ali a gente faz uma integração estudante-empresa, que é até uma das funções do IEL, que é um dos órgãos da CNI. E você vai fazer o quê? Você vai segurar essa mão de obra e esses cérebros no mercado, que é outro problema que nós temos na área de tecnologia. Então nós estamos trazendo uma gama de sugestões para incentivar esse fomento, essa inovação, essa questão mesmo da formação e garantir que esses profissionais fiquem no país.

Outra questão, além dessas exceções e tudo o que a gente já colocou, é o problema da regulação baseada em risco. Entendemos que é muito lógico, muito proporcional, mas o risco precisa ficar muito claro. A gente também entende que o projeto não está muito claro quanto a esses riscos. Todos os desenvolvedores, todos a quem a gente pergunta, têm uma série de dúvidas. Até a mesma sugestão que foi colocada, acho que duas vezes, de trazer também o que não é risco, que está na União Europeia, no art. 6º, a gente também acha que é muito interessante, produtiva como sugestão de texto.

E também, o que não é alto risco não precisa estar aqui dentro, que é a maioria, nós entendemos que o projeto deve ser baseado exatamente no que é o alto risco. E o alto risco deve ser previsto como obrigações de governanças mais gerais – até o Senador Astronauta colocou isso na proposta dele, que também é muito interessante –, técnicas e gerais, que não sejam obrigações de governanças que vão poder ser facilmente superadas num futuro próximo, mas que sejam gerais, porque essa tecnologia é muito granular: vai depender do uso, vai depender do tipo, vai depender do tamanho da indústria, vai depender de quem está desenvolvendo. Essa granularidade deve ser deixada para a regulação setorial. E aí também alguém colocou que o regulador também vai aprender.

Então, a autorregulação regulada, e aí toda a doutrina coloca que o melhor dela é que você faz essa ponte entre regulado e regulador. Você tem uma conversa praticamente, não é diária, mas sempre que forem aumentando e mudando as tecnologias, os regulados trazem as questões, os reguladores colocam as preocupações, e ali você vai ajustando caso a caso, caso e modelo, e a gente vai evoluindo junto, a regulação com a tecnologia. Então a gente entende que essa ferramenta deve estar mais explícita no projeto, nos casos em que ela vai incidir.

Também outra questão que nos causa muita espécie é: se nós estamos partindo para um modelo de sistema de risco, inclusive com listas, a gente entende que – se for partir para esse modelo mesmo, que está sendo avaliado – essas listas deviam ser exaustivas. E aí o que acontece? Por exemplo, avaliações preliminares, que é outra coisa que dá muito problema: devem ser avaliados aqueles casos em que há alguma dúvida. Nos casos em que não se tem nenhuma dúvida, você tem uma lista dizendo o que é para todo mundo fazer e terá uma obrigação regulatória quanto a isso.

Mas o tempo está acabando. Eu vou fazer só a última coisinha. Quanto à questão que também a CNI colocou muito, da sobreposição regulatória, preocupação dessas sobreposições regulatórias entre uma agência, uma autoridade setorial e uma autoridade central, nós entendemos que isso é resolvido com o princípio da especialidade, que tem nos Estados Unidos, onde também tem agências setoriais: deve ficar para aquela agência ou aquela autoridade setorial que tem maior especialidade na matéria para regular aqueles setores não regulados, ou uma capacidade técnica.



SENADO FEDERAL  
Secretaria-Geral da Mesa

Uma autoridade só que vai regular todo mundo: é muito difícil que ela tenha mãos para conseguir fazer isso, dinheiro – e eu acho que feriria o princípio da eficiência da administração e da economicidade, que é você alocar os recursos bem feitos –, e, depois, nós podemos ter muitos problemas concorrenciais. Os setores regulados não estão gostando dessa situação. Por quê? Porque você vai ter uma regulação para soluções iguais no final das contas; cada um com uma regulação por um órgão; ou regulações difíceis, diferentes em que nós teremos problemas de competitividade.

Então, depois, nas perguntas, a gente continua.

**O SR. PRESIDENTE** (Chico Rodrigues. Bloco Parlamentar da Resistência Democrática/PSB - RR) – Nós aqui assistimos atentamente à apresentação de todos os convidados.

Como se vê, esta audiência pública, que é promovida para essas questões mais fundamentais aqui no Senado Federal, é uma forma – é um tempero – exatamente de ouvir as ideias dos vários especialistas, que vão subsidiando a Comissão legislativa do Senado para que possa chegar a uma conclusão e apresentar um projeto pronto e acabado, que venha a ser votado, obviamente – para que um projeto de lei dessa envergadura possa, na verdade, ser utilizado por todo o país. Essa é a função das audiências públicas qualificadas, como esta que o Senador Astronauta Marcos Pontes convidou, para que pudéssemos estar nesta tarde de hoje aqui, nesse debate.

Eu gostaria, até por uma questão de compromisso do Senado com o e-Cidadania, de apresentar aqui – foram muitas perguntas dos internautas – algumas dessas perguntas, e, posteriormente, a Comissão Temporária Interna de Inteligência Artificial ficará responsável pelas suas respostas.

O Gabriel, do Acre, pergunta: "Como o Congresso apoiará a capacitação em IA para que os brasileiros dominem essa tecnologia e posicionem o Brasil como líder global?"

A Danielly, de Rondônia, pergunta: "Como o PL 2.338/2023 assegura que a autorregulação e as boas práticas na [...] [inteligência artificial] sejam eficazes para garantir a ética e a responsabilidade no uso?"

O Dionísio, do Rio Grande do Sul: "Quais metodologias podem ser adotadas para combater o uso abusivo de IAs no contexto de produções científicas/acadêmicas?"

O Jackson, de Rondônia: "De que forma a autorregulação no uso da IA pode contribuir para a proteção de dados pessoais e a privacidade dos usuários no Brasil?"

A Ana, de Minas Gerais: "O que será feito para coibir o uso de [...] [IA] como ferramenta para manipular informações [...] [criando] imagens e áudios falsos?"

Isso, na verdade, hoje, é recorrente, como todos sabem.

A Ana, de São Paulo, pergunta: "Mulheres sul-coreanas estão sofrendo com vídeos *deepfakes* adultos [...] criados [...] [criados por IA sem consentimento]. Como isso será evitado aqui?"

É uma pergunta com resposta? Não sei.

A Isabella, de Mato Grosso do Sul: "Quais medidas serão tomadas para garantir que o uso de IAs não [...] [interfira] na propriedade intelectual de terceiros?"

O Lucas, de São Paulo: "Visto que os algoritmos de inteligências artificiais são [...] [altamente] mutáveis e variados, como o PL prevê lidar com essa [...] [rápida] adaptação?"

E, por fim, o João, do Distrito Federal: "Há garantias para que a IA seja tratada como ferramenta, ao invés de ser considerada como mais uma mão de obra [...] [substituta de] trabalhadores?". (Pausa.)



SENADO FEDERAL  
Secretaria-Geral da Mesa

Agora, nas considerações finais, eu não sei se alguém, algum dos convidados, gostaria de fazer algum comentário em relação a alguma dessas questões. Se puderem fazer, têm o tempo de dois minutos.

*(Intervenções fora do microfone.)*

**O SR. PRESIDENTE** (Chico Rodrigues. Bloco Parlamentar da Resistência Democrática/PSB - RR) – Luis Fernando? Ah, a Christina Aires quer fazer um comentário.

**A SRA. CHRISTINA AIRES CORREA LIMA DE SIQUERIA DIAS** (Para expor.) – Eu acho as perguntas muito pertinentes e eu vou tentar passar rapidamente por cada uma delas.

Acredito que a do Gabriel, sobre capacitação em IA, é uma das preocupações fundamentais da CNI. Ela conversa com a última, do João, do Distrito Federal, sobre como não vai substituir a mão de obra. Nós acreditamos que a Constituição, quando prevê a proteção do trabalhador contra a mecanização – e hoje a gente não é mais máquina, são sistemas... Eu acho que a maior proteção é o treinamento profissional, a formação profissional, para que esses trabalhadores possam ser realocados e saibam usar essa tecnologia, essas ferramentas, que serão mais uma ferramenta de trabalho. Eu ouvi alguém falar que, hoje em dia, os trabalhadores não vão ser substituídos pela IA, mas por aqueles que saibam usar IA. Então acho que essa é uma preocupação que tem que estar mesmo na mira do projeto.

A questão da autorregulação e boas práticas também conversa com essa da proteção de dados e metodologias de usos abusivos.

*(Soa a campanha.)*

**A SRA. CHRISTINA AIRES CORREA LIMA DE SIQUERIA DIAS** – Porque eu acho que, na autorregulamentação, o regulador vai estar muito mais próximo das aplicações e dos usos – é o que a gente coloca. A ferramenta, o sistema não é bom ou mau, o uso dele é que é ruim, que pode ser ruim e que deve ser proibido.

E nessa proximidade entre regulador, entre... como serão desenvolvidos os sistemas, eles vão conseguir talvez combater muito melhor do que com uma lei que não é flexível, que fica parada num momento e que não acompanha a evolução dessas tecnologias. Então, só tentando responder a tudo de uma vez.

**O SR. PRESIDENTE** (Chico Rodrigues. Bloco Parlamentar da Resistência Democrática/PSB - RR) – Eu quero agradecer à Christina Aires pela rápida abordagem dela em relação a algumas perguntas dos internautas... Já ia falando "do astronauta", porque ele não está presente, mas, enfim... dos internautas, que se posicionam sempre em relação aos debates aqui no Senado, o que é muito rico, porque é a curiosidade em relação aos temas que são tratados aqui nos projetos do Senado.

E o Luis Fernando gostaria também de fazer um breve comentário sobre algumas dessas questões.

Dois minutos. V. Sa. tem a palavra.

**O SR. LUIS FERNANDO PRADO** (Para expor.) – Obrigado, Senador, pela fala.

Eu achei curiosas as perguntas porque elas remetem a temas, inclusive, que já são disciplinados hoje na legislação brasileira e já são inclusive proibidos. Então ninguém pode usar a imagem de outra pessoa para fazer *deepfake*, para, de alguma forma, tentar se passar por ela, etc. Inclusive esses usos podem ser feitos sem ferramentas de IA. A gente tem *softwares* de edição de imagem já mais antigos que não dependem de IA para funcionar, e isso já é proibido hoje no Brasil.





SENADO FEDERAL  
Secretaria-Geral da Mesa

Eu não estou falando aqui menosprezando o potencial que a IA pode ter nas relações sociais. Eu acho que vale a gente refletir, sim, sobre o tema, mas talvez regulando usos proibidos em vez de tratar a tecnologia transversal como um todo, como algo que a gente já parte do pressuposto de que ali tem um risco, de que ali tem de fato alguma coisa que precisa ser inviabilizada. A gente tem exemplos de regulações mundo afora que estão apostando muito mais – em vez de criar uma lei geral sobre o tema – em você pegar assuntos que de fato precisam ser modernizados, regulados e aprimorados por causa da IA e regular esses usos, regular essas consequências.

*(Soa a campainha.)*

**O SR. LUIS FERNANDO PRADO** – Então, achei bem interessante a pergunta dos internautas porque nos deu essa oportunidade de mostrar que o PL 2.338 sequer endereçaria a muitos dos pontos aqui, porque eles já estão previstos na legislação brasileira, e não evitaria práticas como essa, porque isso já é vedado. As pessoas não fazem isso porque não têm uma lei de IA, as pessoas não deveriam fazer isso porque a gente já tem hoje Código Civil, a gente já tem várias disposições, inclusive na Constituição, que protegem o direito de imagem, o direito de intimidade das pessoas.

Então, achei excelente o ponto porque nos deixa aqui com esta reflexão final, pelo menos do meu lado: para que caminho a gente está indo? De regular a tecnologia como um todo de maneira transversal? Ou vamos apostar e focar de fato nos usos que devem ser coibidos?

**O SR. PRESIDENTE** (Chico Rodrigues. Bloco Parlamentar da Resistência Democrática/PSB - RR) – Muito bem e muito obrigado, Luis Fernando, pela sua resposta explicativa para os internautas.

E gostaria também de consultar os americanos, a Courtney Lang e o Matthew Reisman, se gostariam também de comentar alguma dessas perguntas. *(Pausa.)*

*Estão online?*

**O SR. MATTHEW REISMAN** (Para expor. *Por videoconferência. Tradução simultânea.*) – Bem, muito obrigado pela oportunidade de responder às perguntas.

Eu acho que a pergunta enfrentou questões importantes que precisamos considerar em relação à regulação, e algumas dessas questões já foram enfrentadas aqui na legislação. Mas o que eu gostaria de dizer é que, por um lado, seria enfrentar isso por meio de diferentes aspectos específicos da legislação, mas também pelo processo de avaliação de riscos.

Os empregadores e desenvolvedores de inteligência artificial vão capturar essas preocupações durante os processos de concepção, e, depois que esse processo, esse projeto for finalizado, será importante a indústria analisá-lo e trabalhar junto com o Governo para identificar boas práticas e fazer a verificação de tudo isso. Nós temos várias experiências bem-sucedidas aqui nos Estados Unidos, que trabalharam nesse desenvolvimento, que ajudaram no desenvolvimento de inteligência artificial aqui nos Estados Unidos e mitigando os riscos. E essa é, talvez, uma oportunidade ou um modelo que também possa se seguir aqui no Brasil depois que a legislação for aprovada.

**O SR. PRESIDENTE** (Chico Rodrigues. Bloco Parlamentar da Resistência Democrática/PSB - RR) – Quero agradecer a participação do Matthew Reisman.

Não sei se a Courtney Lang gostaria também de fazer o seu comentário. *Está online?*  
*(Pausa.)*



SENADO FEDERAL  
Secretaria-Geral da Mesa

**A SRA. COURTNEY LANG** (Para expor. *Por videoconferência. Tradução simultânea.*) – Sim, certamente foram várias questões apresentadas. São questões muito complicadas quando nós começamos a falar na abordagem dessas questões, mas há questões que...

Se de fato essas questões práticas foram aplicadas, que são focadas na mitigação desses problemas, na aplicação de inteligência artificial, no ciclo de aplicação da inteligência artificial... E enfrentar essas questões quando a inteligência artificial entrar em produção, como eu mencionei aqui na geração de conteúdo...

Existe um importante papel a ser desempenhado aqui pelo Congresso: qual é a melhor abordagem, como nós podemos estabelecer a melhor abordagem, o melhor conteúdo e como nós podemos estabelecer diversos desses mecanismos de aplicação e de uso dessas ferramentas em relação a conteúdo gerado pela inteligência artificial. Essas são áreas em que nós precisamos de mais desenvolvimento de pesquisa, e o Brasil pode ampliar essa investigação, como apoiar o desenvolvimento de tecnologia e de pesquisa, por exemplo, a colocação de marca d'água em todos esses conteúdos que são produzidos pela inteligência artificial e a aplicação dessas marcas em todos os conteúdos que são produzidos pela tecnologia de inteligência artificial. Aí nós poderemos chegar a esse potencial uso, esse uso oculto ou malicioso, desse conteúdo.

Tudo isso é importante.

Nós agradecemos a oportunidade aqui que nos deu de participar de tudo isso e estou ansiosa em poder trabalhar, continuar trabalhando com vocês nesse processo de desenvolvimento da tecnologia. Estou aberta para todas essas consultas daqui para frente.

**O SR. PRESIDENTE** (Chico Rodrigues. Bloco Parlamentar da Resistência Democrática/PSB - RR) – Gostaria de agradecer a participação brilhante da Courtney Lang e dizer, na verdade, que os seus comentários finais enriquecem inclusive o debate.

Eu passo a palavra para o Ronaldo Lemos, Diretor do ITS Rio.

**O SR. RONALDO LEMOS** (Para expor. *Por videoconferência.*) – Muito obrigado, Senador.

Tem uma pergunta que eu acho que vale a pena responder, que é a preocupação com a geração de *deepfakes*, de imagens que usam inteligência artificial para criar, muitas vezes, conteúdo pornográfico e conteúdo que afete o bem-estar. A gente tem visto casos em escolas no Brasil de colegas que criam imagens dos colegas em situação vexatória e assim por diante.

Felizmente, essa questão foi endereçada pelo marco civil da internet. O marco civil tem uma provisão sobre imagens sexuais que são divulgadas sem consentimento. Não importa se a imagem é criada por inteligência artificial ou se ela foi filmada do mundo real, mas, em havendo esse tipo de imagem, os provedores são obrigados a remover esses conteúdos imediatamente. É uma medida legal muito eficaz que a lei já endereçou.

Felizmente, a legislação brasileira avançou muito nesse sentido. A gente tem legislações recentes sobre *bullying*, legislações recentes sobre o fenômeno de *stalking* e outras questões. Então, nesse ponto, a legislação está bem avançada. Inclusive, caso aconteça algo semelhante, eu recomendo sempre procurar uma delegacia de crimes cibernéticos, quando elas estiverem disponíveis na cidade em que a vítima habita. Se não estiver disponível, procure uma delegacia normal, mas felizmente, graças à atuação do Congresso Nacional, nosso país está muito avançado com relação a essas questões.

Obrigado.

**O SR. PRESIDENTE** (Chico Rodrigues. Bloco Parlamentar da Resistência Democrática/PSB - RR) – Eu quero agradecer aos convidados, repetindo os nomes: à Courtney Lang, nos Estados Unidos; ao Luis Fernando Prado, que presencialmente, como membro da



SENADO FEDERAL  
Secretaria-Geral da Mesa

Faculdade de Direito de Vitória, atendeu ao chamamento; ao Matthew Reisman, também nos Estados Unidos; ao Ronaldo Lemos, do ITS; e à Christina Aires, da CNI.

Quero dizer que a colaboração de vocês nesta audiência pública foi fundamental. Os postulados que regem o Congresso Nacional, especificamente o Senado da República, são no sentido de que haja participação num debate profícuo, técnico, para que na verdade venha subsidiar esses projetos que são de interesse nacional. E, aqui, o que nós vimos foi exatamente isso, técnicos de altíssimo nível, de qualificação, com contribuições muito proveitosas e, com certeza, a equipe dos consultores irá aperfeiçoar mais ainda esse projeto de autoria do Presidente da Casa, o Presidente Rodrigo Pacheco.

Em nome do Vice-Presidente da Comissão, Senador Astronauta Marcos Cesar Pontes, e em meu nome, queremos agradecer a presença de todos vocês que foram convidados e, também, aos presentes.

Nada mais havendo a tratar, declaro encerrada a presente reunião. *(Palmas.)*

*(Iniciada às 14 horas e 42 minutos, a reunião é encerrada às 16 horas e 07 minutos.)*