



SENADO FEDERAL
Secretaria-Geral da Mesa

ATA DA 17ª REUNIÃO DA COMISSÃO TEMPORÁRIA INTERNA SOBRE INTELIGÊNCIA ARTIFICIAL NO BRASIL DA 2ª SESSÃO LEGISLATIVA ORDINÁRIA DA 57ª LEGISLATURA, REALIZADA EM 01 DE JULHO DE 2024, SEGUNDA-FEIRA, NO SENADO FEDERAL, ANEXO II, ALA SENADOR ALEXANDRE COSTA, PLENÁRIO Nº 7.

Às quatorze horas e quarenta e dois minutos do dia hum de julho de dois mil e vinte e quatro, no Anexo II, Ala Senador Alexandre Costa, Plenário nº 7, sob a Presidência do Senador Astronauta Marcos Pontes, reúne-se a Comissão Temporária Interna sobre Inteligência Artificial no Brasil com a presença do Senador Styvenson Valentim, e ainda do Senador Paulo Paim, não-membro da comissão. Deixam de comparecer os Senadores Carlos Viana, Veneziano Vital do Rêgo, André Amaral, Weverton, Daniella Ribeiro, Vanderlan Cardoso, Nelsinho Trad, Fabiano Contarato, Chico Rodrigues, Eduardo Gomes e Laércio Oliveira. 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 do Senador Astronauta Marcos Pontes (PL/SP), com a finalidade de debater a Categorização e Avaliação de Riscos em Sistemas de IA (avaliação preliminar de riscos; definição e regulação de risco excessivo e alto risco; critérios de classificação de riscos em sistemas de IA), com o objetivo de instruir o PL 2338/2023, que dispõe sobre o uso da Inteligência Artificial, com a participação de Marcelo Almeida, Diretor de Relações Institucionais e Governamentais da Associação Brasileira das Empresas de Software (ABES); Ana Bialer, Membro do Conselho Nacional de Proteção de Dados (CNPd) e Consultora do Movimento Brasil Competitivo (MBC); Mateus Costa-Ribeiro, Fundador da Startup Talisman IA; Felipe França, Diretor-Executivo do Conselho Digital; Patricia Peck, Membro do Comitê Nacional de Cibersegurança (CNCiber); Jefferson de Oliveira Gomes, Diretor de Tecnologia e Inovação da Confederação Nacional da Indústria (CNI); e Rodrigo Scotti, Presidente do Conselho de Administração da Associação Brasileira de Inteligência Artificial (Abria). Nada mais havendo a tratar, encerra-se a reunião às dezessete horas e oito 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:

<http://www12.senado.leg.br/multimidia/eventos/2024/07/01>

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP. Fala da Presidência.) – Declaro aberta a 17ª Reunião da Comissão Temporária Interna sobre Inteligência Artificial no Brasil, criada pelo Requerimento nº 722, de 2023, com a finalidade de examinar projetos concernentes ao relatório final aprovado pela Comissão de Juristas, responsável por subsidiar a



SENADO FEDERAL
Secretaria-Geral da Mesa

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.)*

Aprovada.

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

Esta reunião se destina à realização de audiência pública com o objetivo de debater a categorização e avaliação de riscos em sistemas de inteligência artificial, análise preliminar de riscos, definição e regulação de risco excessivo e alto risco, critérios de classificação de riscos em sistemas de IA, em cumprimento ao Requerimento nº 2, de 2024, da Comissão, da minha autoria.

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

Encontram-se presentes no Plenário da Comissão: Mateus Costa-Ribeiro, fundador da *startup* Talisman IA...

Acho que temos locais aqui, não é?

Mas não são todos.

O.k! *(Pausa.)*

Ah, você vai chamar?

Ele vai chamar no momento correto.

Jefferson de Oliveira Gomes, Diretor de Tecnologia e Inovação do Serviço Nacional de Aprendizagem Industrial (Senai), representando a Confederação Nacional da Indústria (CNI); Marcelo Almeida, Diretor de Relações Institucionais e Governamentais da Associação Brasileira das Empresas de *Software* (Abes); Felipe França, Diretor-Executivo do Conselho Digital; e Rodrigo Scotti, Presidente do Conselho de Administração da Associação Brasileira de Inteligência Artificial.

Encontram-se também presentes, por meio do sistema de videoconferência, Patricia Peck, Membro do Comitê Nacional de Cibersegurança, e Ana Bialer, Membro do Conselho Nacional de Proteção de Dados e Consultora do Movimento Brasil Competitivo.

Eu agradeço a presença de todos. Nós ainda estamos esperando mais um participante se confirmar neste último momento, mas eu começo agradecendo a presença de todos aqui. Nós temos também a presença conosco do ex-Senador João Costa, do Tocantins. Obrigado pela presença, que certamente nos enriquece com seu conhecimento e presença aqui também.

Gostaria de aproveitar e agradecer a presença de todos que se encontram aqui presencialmente e também daqueles que nos acompanham através da TV Senado e das redes do Senado, assim como dos participantes que estão no *online* aqui conosco.

Esta audiência pública tem uma finalidade bastante específica, como foi lido. Mas, de forma geral, o que acontece? Nós tivemos a proposição de vários projetos de lei, inclusive na Câmara dos Deputados, como o Projeto nº 21, que foi relatado pela Deputada Luísa Canziani, e outros projetos também. Depois, em 2022, veio a Comissão de Juristas com o 2.338. Esse 2.338 foi extremamente discutido aqui, através desta Comissão, e eu presidi a maior parte das audiências como esta ao longo do período.

Com base nos resultados dessas audiências públicas, nós fizemos, a muitas mãos, uma emenda, que é a Emenda nº 1. Para quem tiver curiosidade, entre no Portal do Senado, no PL 2.338, de 2022, e



SENADO FEDERAL
Secretaria-Geral da Mesa

encontrará lá a Emenda nº 1, que foi a proposição feita através dos resultados dessas audiências públicas.

É uma proposição mais principiológica, obviamente, da maneira como eu acredito que tem que ser: principiológica, para que possa acompanhar o desenvolvimento da tecnologia sem ficar obsoleta no dia seguinte do seu lançamento, e que propõe a criação de um comitê formado por representantes de agências reguladoras, da sociedade civil, do setor produtivo, da academia, da comunidade científica. De forma que esse conselho, aí sim, faça a atualização periódica das tabelas de risco, probabilidade e impacto de riscos, porque isso, obviamente, se altera também.

Você não consegue, de maneira nenhuma, regular uma tecnologia, mas você pode e deve regular os resultados, as implicações da tecnologia que são traduzidas através de riscos. Portanto, é necessário esse trabalho contínuo, da mesma forma assim como é feito com a proteção de dados, e a proposta nossa de que seja feita a coordenação pela Autoridade de Proteção de Dados, que já tem um conselho semelhante...

Na verdade, eu gostaria de ver no Brasil três conselhos ali: proteção de dados, inteligência artificial e segurança cibernética, espalhados com representantes nos diversos setores, de forma que a gente tenha uma aplicação dos princípios da lei dentro de cada setor, sem haver conflito entre a lei e as regulações de cada setor, ou seja, sem haver interferência. E que possa ser rapidamente adaptado para as necessidades que existam e vão aparecer com o desenvolvimento da tecnologia. Algo que é de alto risco hoje, pode ser de baixo risco amanhã; alguma coisa que é de baixo risco hoje, pode ser de alto risco amanhã. Então, é importante sempre se ter isso em mente.

Nós tivemos o 2.338, tivemos a Emenda nº 1, e depois o nosso Relator, que é o Senador Eduardo Gomes, veio com um texto intermediário – que é o texto que é razão desta audiência –, que foi apresentado. Ele não representa exatamente a 2.338; ele tem melhorias, por exemplo, com relação a fomento, o que é importante – inclusive está melhor do que a Emenda 1, com relação a fomento para o desenvolvimento no Brasil, *sandboxes* e assim por diante –, mas, como ele ainda tem muita diferença com a emenda e com o 2.338, eu sugeri que nós tivéssemos audiências públicas para, de novo, trazer a sociedade para discutir, porque isso não é uma coisa que, na minha opinião, tem uma pressa gigantesca para ser feita – eu já coloco assim de cara – e também é importante a participação de todo mundo para discutir, porque esta aqui é uma Casa que representa a população. A gente tem que ter a participação de todos. O que não pode acontecer é chegar ao final e alguém falar assim: "Eu não fui ouvido, eu não pude participar", e assim por diante. Então, a gente abre justamente a audiência pública para que isso seja possível. Num país democrático, num regime democrático, é extremamente importante essa participação. Vale lembrar que inteligência artificial interfere na vida de todo mundo, vai interferir cada vez mais, positiva ou negativamente, então é importante que todo mundo participe. Isso não é uma legislação qualquer.

Agora, acontecem as acelerações ou as pressões aqui para que as coisas andem rápido. A ideia – já para colocar aqui –, inicialmente, era que nós tivéssemos cinco audiências públicas, como eu propus. Isso foi discutido. E foi solicitado pelo nosso Presidente da Comissão, que é o Senador Carlos Viana, que nós reduzíssemos de cinco para três audiências públicas, para acelerar o processo. Obviamente, ele tem as suas outras necessidades ou pressões, então passam a ser três audiências públicas, a pedido do Presidente da Comissão, que vão ser feitas esta semana, e existe a possibilidade, inclusive, de haver a votação, nesta semana, do texto.

O Senador Eduardo Gomes vai trabalhar em cima desse texto. Eu espero – eu não posso interferir, obviamente; ele é o Relator, eu não posso obrigá-lo a colocar nada ali – que ele receba as informações



SENADO FEDERAL
Secretaria-Geral da Mesa

que nós estamos discutindo, ou que vamos discutir aqui e nos outros dias, e faça as devidas adequações ao texto, para que essa legislação atenda aos seguintes critérios, que eu vejo como essenciais: primeiro, para que ela favoreça o fomento para o desenvolvimento dessa tecnologia no Brasil; segundo, para que ela possa permitir o desenvolvimento da utilização ou aplicação de inteligência artificial nos diversos setores do Brasil sem causar restrições, que seriam extremamente danosas à competitividade do Brasil em todos os setores. E, quando falo "todos", estou falando de todos os setores, mesmo. Isso aí é uma tecnologia que interfere em todos os setores, então eu espero que ele tenha sensibilidade e que ouça os apelos dos diversos setores.

Eu estou aqui como uma ferramenta para trazer isso, através das audiências. Eu espero que cada representante apresente as suas propostas também, para que as propostas fiquem claras, publicamente claras, para que cheguem ao Relator de forma pública também. Ele obviamente tem recebido bastante de outros *inputs*.

Em terceiro lugar, que essa legislação seja suficientemente desenvolvida e flexível para se adaptar às variações da tecnologia, o que implica aquilo que eu falei do comitê, etc., que possa fazer a atualização frequente dos parâmetros de risco, que a gente vai discutir aqui hoje. Imagino que quem está aqui já trabalha com risco há bastante tempo; eu também, foram mais de 30 anos fazendo investigação e prevenção de acidentes. A gente sabe muito bem que a coisa importante com relação a risco é você notar a situação, o impacto e a probabilidade e, com isso aí, você ter critérios objetivos para distribuir, ou para classificar os riscos e depois tomar as providências adequadas, como mitigação, eliminação, transferência, o que for, mas precisa ter classificação objetiva. Esse é outro critério que eu gostaria de ver nessa regulação, que são critérios objetivos e não subjetivos, o que levaria a uma margem de insegurança jurídica, o que não é adequado para o país, muito menos para o investimento no país.

Então, são esses alguns pontos que eu gostaria de ver nessa lei.

Eu espero que o nosso Senador Eduardo Gomes, Relator, esteja atento a tudo isso e que ele nos traga um texto adequado para ser votado. Se tiver um texto bom, excelente. Votar o mais rápido possível é bom, sem dúvida nenhuma, mas, novamente, friso que eu – eu, particularmente, Senador Astronauta Marcos Pontes, uma pessoa, um Senador entre 81, falando aqui – não vejo razão de pressa, mesmo porque observar o que é feito em outros países pode ser até benéfico para que você corrija os erros que eles fizeram, ou evite os erros que eles fizeram, e copie aquilo que funciona, adeque para as nossas condições.

Bom, dito tudo isso – desculpem aí a longa introdução –, também uma informação que eu recebi hoje de um meio de imprensa dizia que alguns dos nomes estavam incorretos com o cargo, vamos dizer assim, o tipo de organização de que participavam. Então, às vezes acontecem algumas trocas, mas eu vou pedir, depois, para o pessoal confirmar realmente: "Olha, eu faço parte de tal", para não ter esse erro.

A ideia geral... Eu não tenho muita preocupação exatamente com isso, porque a ideia geral, como eu falei no início, é que seja democrático, que todos possam participar. Então, que a gente tenha representantes da academia, comunidade científica, setor produtivo... que todo mundo possa participar. Essa é uma audiência pública, paga com dinheiro público, para o bem do público, e a gente precisa fazer isso acontecer aqui.

Vamos começar, então, com como vai funcionar o sistema.

Eu já tenho algumas perguntas e comentários aqui, através do e-Cidadania. Eu vou pedir para distribuí-las aos participantes, porque no final a gente vai responder essas perguntas ou os comentários.



SENADO FEDERAL
Secretaria-Geral da Mesa

A dinâmica será: eu vou passar a palavra, na sequência, para os apresentadores. Vocês têm dez minutos para fazer a apresentação. Para quem está aqui é fácil, porque tem um relógio na parede, você vê e vai tocar algo como isto...

(Soa a campainha.)

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – ... quando estiver faltando um minuto. E, depois, aos 15 segundos, vai entrar uma voz feminina, muito convincente, para falar que faltam 15 segundos. E aí, terminam a ideia dentro desse tempo. Então, vamos dizer, gerenciem o tempo por aqui.

Para aqueles que estão no remoto – tem três pessoas no remoto –, eu vou pedir para prestarem atenção por lá, controlarem o tempo por lá, porque não vão ter essas dicas daqui, para a gente não ultrapassar muito o tempo.

O.k. Então, vamos começar com uma participação remota, da Patricia Peck. Ela é membro do Comitê Nacional de Segurança Cibernética e vai participar remotamente.

Está tudo conectado já? *(Pausa.)*

Eu peço então para passar a palavra à Patricia Peck, que tem dez minutos para sua apresentação. Agora, estou vendo-a aqui no monitor também.

Patricia, tudo bem? Obrigado pela participação. Então, dez minutos para sua apresentação. Controle o seu tempo por lá. Obrigado, um abraço. Está contigo.

A SRA. PATRICIA PECK (Para expor. *Por videoconferência.*) – Boa tarde a todos.

Agradeço a oportunidade de estar presente aqui numa das audiências públicas, conforme mencionado pelo Senador Marcos Pontes. Também aproveito para cumprimentar todos os colegas que irão apresentar painel nesta data e cumprimentar todos os Senadores, em especial o Senador Marcos Pontes, o Senador Eduardo Gomes e o Senador Carlos Viana, pelo trabalho que vem sendo realizado nessa matéria.

Eu compartilhei aqui a minha tela com o objetivo desta audiência, tendo sido feito o convite para tratar especificamente da questão da classificação e categorização de riscos. E, já estando atuando há mais de 20 anos no Brasil com direito digital, com cibersegurança, proteção de dados, inclusive aqui também representando o Instituto Noberto Bobbio e o Instituto de Cidadania (IPCD), que são entidades que dialogam diretamente com vários setores produtivos, em especial com o setor de ensino, eu queria destacar da fala do Senador na sua abertura que o objetivo de todos nós é viabilizar um marco legal brasileiro que possa estimular o uso, o fomento da tecnologia da inteligência artificial no Brasil. Da forma como apresentado e proposto, seguindo um pouco a linha da diretriz europeia, toda a redação começa... ou seja, a forma de lidar com os próprios requisitos apresentados pelo marco legal envolve sempre o enquadramento na classificação de riscos.

Tendo em vista já o exaustivo trabalho realizado por esta Comissão – muito bem-feito –, eu vou dar alguns exemplos de situações concretas que podem acontecer, que depois podem ser comentadas pelos colegas de painel.

A primeira questão envolve a avaliação preliminar trazida pelo art. 12, que deve ser realizada inicialmente e que pode ser feita tanto pelo desenvolvedor como por aquele que vai fazer uso da inteligência artificial. O primeiro desafio é o enquadramento adequado na classificação de riscos, por isso ela traz alguns exemplos. Entre esses exemplos, alguns podem, de algum modo, comprometer um pouco a realidade atual que nós temos no Brasil, principalmente aqui voltada a alguns tipos de perfis, como a importância de se aumentar a segurança e protocolos de segurança em entidades de ensino.



SENADO FEDERAL
Secretaria-Geral da Mesa

O segundo ponto é o risco de sofrer uma reclassificação. Por mais que haja previsão ali do contraditório e da ampla defesa, sabemos que os investimentos em inteligência artificial são altos e são de médio para longo prazo. Então, uma mudança de reclassificação de risco pode significar não apenas a necessidade de mais controles e de implementar governança, o que traz, com toda certeza, custos, mas também até de se reclassificar para um risco inaceitável, o que gera, sim, uma certa insegurança.

Terceiro, o desafio de se trazer uma lista de tipos de sistemas de IA vedados. Como muito bem colocado pelo Senador, uma coisa é tratarmos risco, outra é exemplificarmos, na própria legislação, tipos de tecnologias que podem, de algum modo, na sua aplicação, ter impactos distintos, daí a importância também da revisão.

E, quarto, redação muito aberta. Só para dar um exemplo, no art. 13, em alguns dos seus incisos, há a possibilidade ali de um entendimento de que, de alguma maneira, cause ou seja provável que cause danos à segurança e a outros direitos fundamentais.

Então, dito isso, eu vou dar os exemplos.

No art. 13, eu quero destacar, especificamente, o inciso VII:

Art. 13. São vedados a implementação e o uso de sistemas de inteligência artificial:

.....
VII - sistemas de identificação biométrica à distância, em tempo real e em espaços acessíveis ao público, com exceção das seguintes hipóteses:

.....
Toda a lista das hipóteses tratadas aqui está muito mais voltada para a utilização em segurança pública. E eu digo isso, então, porque fica um desafio, ou ainda em revisão aqui da matéria apresentada, para quando são usados para fins de segurança privada, como eu vou dar exemplos aqui a seguir, dentro do tempo que se apresenta.

Uma coisa é definir que tem que se garantir ali uma proteção a um risco de viés e a risco de discriminação. Esse ponto está completamente correto dentro da legislação e acompanha todas as tendências internacionais. No entanto, na hora em que se traz ali a exceção mais específica para a segurança pública, isso pode dar margem à dúvida em casos de utilização privada: em seguranças que acontecem em *shopping centers*, em lojas, hoje, no comércio, nas agências de banco, nas escolas, que aumentaram muito o protocolo relacionado à segurança. Eu quero destacar isso porque pode alcançar até o agronegócio brasileiro, que hoje faz uso de *drones* que fazem monitoração de perímetro em produções e também em plantas industriais. E, numa leitura inicial, não se identifica ali claramente de que maneira seria legítima, então, a sua aplicação em segurança privada.

Mesmo fazendo a leitura do *hall* previsto em alto risco, em que se poderia imaginar a possibilidade de enquadrar alguns dos exemplos aqui colocados no art. 14, ainda assim ele foi escrito com uma redação bem mais restritiva, considerando sempre a necessidade de autenticação específica ali daquele usuário. Então, se pensarmos o uso de videovigilância com reconhecimento facial em perímetro de acesso público mais amplo – acabei de mencionar também arenas esportivas e clubes –, resta uma dúvida de que não é apenas o momento de autenticação para entrada, mas, sim, do que acontece naquele ambiente, para coibir situações que possam trazer riscos àquela própria comunidade que está naquele local – por isso, alcançar clubes e agremiações.

Além disso, eu queria reforçar que o Brasil vivencia um ambiente em que as próprias entidades de ensino investiram bastante nesses protocolos desde o ano passado. Daí, por isso, na harmonização entre a necessidade de nos protegermos em riscos de segurança de um Brasil em que aumenta a conta da fraude e em que também tem tido episódios de ataques a esse ambiente dessas entidades *versus*



SENADO FEDERAL
Secretaria-Geral da Mesa

um excesso de cautela, lógico, numa primeira redação do marco legal, esse tema seria melhor acomodado de forma setorizada.

Eu quero dizer que, a partir do momento que citamos a tecnologia específica de uso de biometria ou de reconhecimento nesses perímetros em que há público, acabamos depois tendo que tentar listar uma série de exceções, e quem não está na exceção deve resolver como aquilo que já está implantado e aquilo que ainda pode ser implantado.

Então, ou evitamos esse grau de detalhamento, de ficar trazendo exemplos de tecnologia e aplicações na lei, deixando apenas os parâmetros de risco – como bem comentado pelo Senador Marcos Pontes – ou teríamos ainda que acrescentar outras finalidades de uso no art. 13, que não conseguiram ser bem tratadas ainda na redação atual, como quando temos o uso privado em termos de protocolos de segurança e prevenção, cumprimento de contratos e proteção de vulneráveis, quando é especificamente para a sua proteção.

A legislação traz um cuidado adicional, para não haver um uso indevido pela inteligência artificial, naquilo que se refere à criança e ao adolescente, mas existe o uso para a proteção da criança e do adolescente.

Eu queria terminar a minha fala, destacando que é extremamente importante que o marco legal possa trazer quais são os requisitos mínimos de fábrica que têm que ser observados, principalmente no tocante à cibersegurança.

Nos últimos dois anos, tendo aqui feito uso de muitas dessas soluções, dá para se observar que essa padronização, que tem que vir de fábrica, pensando que hoje isso é realizado não só no Brasil, mas em outros países, é essencial ser prevista numa exigência legal.

Então, vemos, de um lado, um detalhamento, talvez aqui para reflexão de todos, excessivo, ao mencionar algumas das situações e algumas das tecnologias, e, por outro lado, ficou bem raso ou ainda não tão bem tratado o que tem que vir de requisito de fábrica, como o exemplo de prevenir riscos de envenenamento de dados, que é um dos riscos mais graves que temos visto com relação a esse tipo de solução.

Queria agradecer a oportunidade da fala aqui, nesta audiência pública, e reforçar a importância desta abertura e deste momento, para que possamos aqui complementar uma visão e evoluir para o marco legal da IA, que seja extremamente assertivo para fomentar o uso da tecnologia, mas não com um impacto demasiado, que possa até afetar aqueles que já estão usando ferramentas que protegem o próprio cidadão.

Muito obrigada.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Obrigado, Patricia.

Parabéns pela apresentação.

Concordo muito com o que você falou, principalmente no tocante a se colocarem muitos detalhes dentro da lei que vão se tornar rapidamente obsoletos e prevenir ou atrapalhar o desenvolvimento e o uso da tecnologia no Brasil. O que a gente não quer é justamente isto: restringir o desenvolvimento da tecnologia aqui no Brasil.

Então, quando você restringe demais, devemos lembrar que, em todo setor de tecnologia – e quero lembrar que eu também era Ministro da Ciência, Tecnologia e Inovação e, portanto, sei muito bem do que eu estou falando –, é muito fácil para uma empresa de tecnologia decidir ficar no Brasil ou ficar fora do Brasil; e depois a gente vai ter que comprar dados lá de fora, vai ter que comprar os *softwares* de fora, vai ter que comprar... E, certamente, o resultado é pior.



SENADO FEDERAL
Secretaria-Geral da Mesa

É mais arriscado não usar a inteligência artificial, desenvolvendo no país, do que usar – com certeza isso é fato.

E também a introdução de Security by Design, vamos dizer assim, colocar segurança cibernética desde o início do desenvolvimento do projeto é essencial. Realmente, essa, sim, é uma parte essencial e é, vamos dizer assim, principiológica, ou seja, ela tem que estar... Se eu penso em um projeto de lei, o que tem que ficar na peça principal? No meu ponto de vista, é coisa que não vai mudar daqui a 5 anos ou 10 anos. O que vai mudar daqui a 5 anos ou 10 anos tem que estar num sistema que possa ser atualizado rapidamente. É por isso que o risco é a maneira correta de se trabalhar isso e fora do sistema como um todo, para não ficar fazendo todo o processo legislativo... Imaginem, uma vez por ano, ter que submeter a mudanças o projeto? A gente nunca vai conseguir ter um projeto de lei viável se isso for feito.

Eu passo agora a palavra ao Mateus Costa-Ribeiro, fundador da *startup* Talisman IA.

Mateus, por favor, dez minutos.

O SR. MATEUS COSTA-RIBEIRO (Para expor.) – Senador Marcos Pontes, boa tarde.

Como eleitor residente de São Paulo, para mim é um orgulho ver a liderança de V. Exa. neste projeto, representando o Estado de São Paulo. Obrigado pelo convite para estar presente hoje e poder compartilhar um pouco sobre a minha experiência na Talisman, de como é construir uma empresa de inteligência artificial no Brasil.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP. *Fora do microfone.*) – Interessante. Obrigado.

O SR. MATEUS COSTA-RIBEIRO – Um membro da Marinha americana disse, certa vez, que ninguém nunca reclamou de um discurso por ele ter sido curto demais. Então, eu vou ser breve. (*Risos.*)

Eu separei a minha fala em três pedaços: a primeira, quem sou eu; a segunda, como o projeto de lei no formato atual vai prejudicar *startups* brasileiras de inteligência artificial; e, em terceiro lugar, eu trouxe algumas sugestões específicas de como ele poderia evoluir.

Primeira parte: quem sou eu. Eu fiz Direito na Universidade de Brasília, a 3km do prédio do Congresso Nacional, onde nós estamos, e aos 18 anos fui o advogado mais jovem do Brasil. Ainda naquela oportunidade, eu fui a pessoa mais jovem também a ganhar uma ação no Supremo Tribunal Federal e, na sequência, fui fazer o meu mestrado na Harvard Law School, ainda na área do Direito. E, naquele momento, Senador, quando eu me formei, eu tive a oportunidade de fazer o New York BAR Exam, que é a prova da OAB americana, e de começar a advogar no Milbank, que é um escritório de advocacia conhecido de Nova Iorque.

Em outras palavras, naquele momento, o sonho americano de que tanto falam, o *american dream*, estava aberto para mim, inclusive com visto e possibilidade de trabalhar fora. Mas a minha reação, naquele momento, foi, na verdade, apostar no sonho brasileiro. Em outras palavras, depois de ter feito o mestrado em Direito em Harvard, como a pessoa mais jovem, e de ter sido reconhecido pelo *The Wall Street Journal* como a pessoa mais jovem também a passar no BAR de Nova Iorque, entre todos os americanos, eu aprendi que o Brasil não deixa a desejar para nenhum país e que nós temos aqui talentos capazes de disputar com talento americano, chinês, etc.

Agora, a diferença entre o Brasil e os Estados Unidos é que, nos Estados Unidos, no momento de dificuldade, as pessoas que estudaram nas melhores universidades, as pessoas mais capazes, enfrentaram a tributação britânica, a escravidão, em vez de pegar as malas e ir para um país mais desenvolvido, que foi a decisão que eu tomei. Eu, então, voltei para o Brasil e fundei, junto com alguns sócios da área de engenharia, a *startup* Talisman, e a gente tem, entre os nossos investidores, o



SENADO FEDERAL
Secretaria-Geral da Mesa

Presidente Global do Google, o John Hennessy, e também o ex-CEO da Thomson Reuters, o Tom Glocer. E eu estou aqui para compartilhar um pouco de como o projeto de lei pode impactar e prejudicar as *startups* brasileiras.

Eu começo, então, a segunda parte da minha fala: como *startups* brasileiras serão prejudicadas pelo PL 2.338.

Hoje, Senador, uma *startup* de inteligência artificial faz uso de modelos proprietários, construídos dentro de casa, e de modelos públicos, tipicamente construídos pelas *big techs* – pela Anthropic, pela Meta, pela OpenAI e pelo Google. Em outras palavras, as *startups* brasileiras passam e consomem esses modelos de inteligência artificial como se fossem um posto de gasolina, e eles são fundamentais para que uma *startup* brasileira, como a nossa, consiga desburocratizar o Brasil, ler documentos com inteligência artificial e acelerar crédito, que são os produtos que nós oferecemos.

E o projeto de lei traz dois grandes problemas: a imposição de que os modelos de inteligência artificial de alto risco passem por uma espécie de censura prévia em que esses modelos precisam ser construídos e desenhados de uma determinada maneira; e, ao mesmo tempo, a criação de uma ameaça de R\$50 milhões ou 2% do faturamento da empresa, o que for maior, que é uma espécie de espada de Dâmocles, ou seja, a lei é abstrata, é ambígua, é principiológica, mas, se descumprida, tem uma espada pronta para cortar a sua cabeça.

Eu acho que, por esses motivos, o Brasil precisa ir a outra direção. Hoje, a cada dez *startups* que surgem no mundo, quatro são de inteligência artificial. Eu acho que a manchete que o PL traz é: "Não venha abrir a sua empresa no Brasil". Isso por três motivos. O primeiro deles é porque o acesso aos melhores modelos será dificultado, e isso já aconteceu antes. Em agosto do ano passado, o Google lançou o modelo Gemini, que foi lançado no Brasil seis meses depois de ser lançado nos Estados Unidos, e a principal competidora da Talisman é a Harvey, uma empresa americana muito bem financiada, com ex-funcionários do Google, da Meta, e que oferece seu serviço para empresas brasileiras, sem nunca ter pago imposto no Brasil, sem contratar ninguém no Brasil e com talentos que não passarão pela regulação que está sendo aprovada aqui...

O mercado de IA é globalizado, e o projeto nada mais faz do que criar barreiras que serão aplicadas só para a inovação brasileira. A minha sensação, Senador, é de que o meu país não admira a minha escolha de ter apostado no sonho brasileiro e de ter feito e começado a empresa no Brasil, que é o que nós vamos continuar fazendo.

O segundo grande motivo que prejudicará as *startups* brasileiras é a capacidade de contratar talento global; 92% dos donos de *startups* dizem, com unanimidade, que a principal dificuldade é contratar talentos.

Hoje, o Brasil tem um déficit de engenharia de *software* de 530 mil engenheiros. Portanto, a gente tem feito um trabalho hercúleo na direção de trazer talentos de ponta. Nós trouxemos, por exemplo, um engenheiro de ponta da Scale AI, que é uma empresa de inteligência artificial do Vale do Silício, trouxemos um engenheiro de uma empresa canadense; todos brasileiros. Nós estamos repatriando as pessoas, e esse projeto vai nos prejudicar a realizar esse movimento. Por quê? Porque ele cria uma dificuldade de um engenheiro agora ter que dedicar, no mínimo, 30% do seu tempo para explicar o que está sendo feito, independentemente do tipo de modelo que está sendo aprovado.

E eu acho que o terceiro ponto que prejudica as *startups* locais é que o PL está na contramão do que os principais países líderes, nessa área, decidiram executar. Hoje, no *site* de imigração do Governo americano, existe uma manchete: "Traga o seu talento de inteligência artificial para os Estados Unidos". Eles criaram o Visto O-1, com toda a razão – foi inteligente da parte deles –, para acelerar a atração de



SENADO FEDERAL
Secretaria-Geral da Mesa

talentos de inteligência artificial, por exemplo, vindos do Brasil. A China fez o mesmo: investiu US\$1 bilhão na MiniMax, que é a versão chinesa da OpenAI. Enquanto isso, nenhuma sinalização foi feita pelo Governo brasileiro.

Em paralelo, um exemplo paradigmático disso é a Hyperplane, uma empresa genial fundada por três brasileiros no Vale do Silício, que pegaram as malas, foram para lá, atendem, de lá, empresas brasileiras. E esse movimento vai acontecer cada vez mais se nós impusermos, de maneira prematura, uma regulação que só existe no Brasil, que não foi espelhada nos países que lideram o desenvolvimento de inteligência artificial, a exemplo dos Estados Unidos, e que, portanto, vai prejudicar o nosso foguete antes mesmo de ele sair da atmosfera – para usar uma analogia com a sua profissão, Senador.

Eu acho que... Eu queria concluir, conectando com a terceira e última parte da minha fala, apontando, de maneira específica, como é que esse tema deveria ser tratado. Eu acredito que a melhor abordagem já existe. Se uma empresa coloca e comercializa um modelo que vai causar prejuízo, o Código Civil obriga essa empresa a reparar os danos que forem causados. Da mesma forma, o Código de Defesa do Consumidor, se essa empresa for Direct to Consumer, se ela vender um produto diretamente para os consumidores, essa empresa inclusive vai ter o ônus da prova invertido, ela vai ser obrigada a indenizar. Existe, em outras palavras, um marco regulatório, da mesma forma que existe regulação para pornografia infantil em meios digitais, da mesma forma que existe regulação de notícias falsas, ou seja, quem usa um computador do jeito errado é punido. Não por isso, existe uma regulação sobre como pesquisadores podem pesquisar ou pensar em qual vai ser a próxima geração de computadores e como ela será construída.

E portanto, para ser mais específico, eu acredito que o art. 12 cria uma posição humilhante para os desenvolvedores brasileiros, na medida em que eles precisam fazer uma avaliação prévia de modelos independentemente do risco.

(Soa a campanha.)

O SR. MATEUS COSTA-RIBEIRO – Então, se você faz um modelo que compara e separa gatos e cachorros, esse modelo vai passar por esse tipo de avaliação prévia. Isso não faz sentido, isso é um modelo inofensivo.

Da mesma forma, o art. 15 abre a porta para o Sistema Nacional de Regulação de Inteligência Artificial (SIA) categorizar qualquer modelo como modelo de alto risco, porque basta que eles atinjam um alto número de pessoas para que ele se encaixe nessa categoria e aí então, passe a ser obrigado a cumprir uma série de regras que só grandes empresas conseguem cumprir. Uma *startup* não consegue cumprir. Em outras palavras, além de ser ambíguo e abrir a porta para o regulador abusar de um poder que está sendo conferido pelo Poder Legislativo, da mesma maneira, o projeto de lei não protege os menores, não protege a competição, na verdade, protege os grandes.

Para concluir, eu acredito que o melhor lugar para a regulação da inteligência artificial seja o tempo, a paciência. O ChatGPT foi lançado há menos de dois anos.

(Soa a campanha.)

O SR. MATEUS COSTA-RIBEIRO – Não existe amadurecimento para o tema, e eu acredito que há um risco grande de as *startups* serem prejudicadas e não apoiadas como é a intenção do Senado brasileiro.

Muito obrigado.



SENADO FEDERAL
Secretaria-Geral da Mesa

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Mateus, primeiro, parabéns. Parabéns pela carreira, parabéns pela brilhante apresentação. Eu achei extremamente precisa no ponto. Concordo em número, gênero, grau, eu falo 100%.

Isso é uma coisa que eu tenho batido na tecla aqui muitas vezes. Todo mundo que está aqui pelo menos já deve ter me ouvido falar um pouco sobre esses temas. Aliás, eu também recebi algumas críticas, falando assim: "Poxa, você está chamando só as empresas grandes para falar." Não, claro que não. Está aqui uma representação de *startup*. É importante. A gente está abrindo para todo mundo.

Devem ter me acompanhado muito tempo lá como Ministro, trabalhando inclusive em programas que deram muito resultado no Brasil, como o IA², Inovação Aberta e Inteligência Artificial, financiado pela Finep, na minha gestão, porque eu acredito que a solução de grande parte dos problemas que a gente tem está nas *startups*, na movimentação rápida das *startups*. A gente tem o marco legal das *startups*, a gente está trabalhando também para facilitar, alinhar com a Lei do Bem também; ou seja, a gente precisa ir nessa direção de favorecer as nossas empresas, especialmente as empresas pequenas, porque elas têm muito mais agilidade, um potencial muito grande de desenvolver muita coisa no Brasil.

Eu acho que esse tipo de legislação, quando é feita baseada no medo, como a gente viu originalmente – eu falei aqui, algumas vezes, sobre isso –, atrapalha e atrapalha muito o desenvolvimento. Atrapalha demais, porque, antes de você ver qualquer tipo de resultado, você já freia tudo.

Eu trabalhei, como eu falei no começo, por 30 anos, com investigação e prevenção de acidentes aeronáuticos. Você quer saber a maneira mais óbvia e mais ineficiente de se terminar com os acidentes? Param-se os voos, põe-se o avião no hangar e ninguém mais voa, certo? Resultado: zero de resultado prático operacional. A gente precisa aprender a trabalhar com o risco, a gente precisa aprender a medir isso aí de forma pragmática, para não deixar interpretações incorretas e não deixar ficar a insegurança jurídica, o que atrapalha o desenvolvimento.

As pessoas vão desenvolver lá fora. Àquelas que vêm, eu agradeço muito por essa... Eu também trago isto aí dentro de mim: depois de tanto tempo trabalhando fora do Brasil, vir para cá, para trabalhar aqui no Brasil, para ajudar o nosso país... Para isso, a gente precisa ajudar essas pessoas que querem desenvolver no Brasil.

Então, são extremamente importantes muitos pontos que você falou. É bom lembrar que tudo que é falado aqui, todas as apresentações ficam gravadas, ficam no portal. Então, para aqueles que estão nos acompanhando: passem, compartilhem essas informações, para que isso seja registrado e compartilhado. Tudo que é audiência pública é com essa finalidade.

Então, parabéns pela apresentação, parabéns pelos pontos que você trouxe, que são extremamente preocupantes mesmo. Já levantei várias vezes aqui: ter predeterminado – vamos chamar assim – o que é alto risco e o que não é, isso não funciona; na minha opinião, não funciona.

Você precisa ter isso aí sendo feito de uma forma dinâmica, acompanhando o desenvolvimento da tecnologia, feito pelos devidos setores, para não ter conflito – ou para reduzir os conflitos, pelo menos –, e que exista mais objetividade em determinar: "Eu vou fazer uma inteligência artificial, vou desenvolver uma inteligência artificial, eu me classifico aqui". Você pode até fazer uma auditoria depois, para saber se aquilo é realmente lá, mas "Eu me classifico aqui e eu tenho como demonstrar aquilo lá". Agora, ficar à mercê de interpretações não é bom, de jeito nenhum.

Então, parabéns por trazer esses pontos. É exatamente essa a ideia de se ter uma audiência pública, trazendo os diversos setores aqui.



SENADO FEDERAL
Secretaria-Geral da Mesa

Eu passo a palavra agora ao Rodrigo Scotti, Presidente do Conselho de Administração da Associação Brasileira de Inteligência Artificial.

Rodrigo, por favor. Dez minutos.

O SR. RODRIGO SCOTTI (Para expor.) – Obrigado a todos e todas.

Agradeço a oportunidade e o espaço, Senador Astronauta Marcos Pontes. É uma honra dividir a mesa com os outros renomados especialistas. Mateus, excelente fala! Muito bom mesmo.

Bem, é uma discussão extremamente fundamental acerca do futuro tecnológico do Brasil, é este o momento em que a gente está aqui, hoje.

Primeiramente, peço a permissão também para me apresentar. Eu sou o Rodrigo Scotti e sou um dos fundadores da Abria (Associação Brasileira de Inteligência Artificial), uma associação sem fins lucrativos e que hoje reúne uma comunidade de mais de 650 empresas de diversos tamanhos, diversos setores, muitas delas *startups*.

E eu não venho aqui hoje somente por conta da Abria, mas também como desenvolvedor de inteligência artificial, desenvolvedor e fundador de uma *startup* nacional de inteligência artificial – IA brasileira. E, principalmente, a gente foca em modelos de linguagem, que são tecnologias que estão ali tocando a IA generativa, que é o que está dando o que falar recentemente no mundo.

Como empreendedor e representando empreendedores do setor, eu gostaria de, primeiro, agradecer a oportunidade de fala pela primeira vez aqui nesta Comissão.

Eu gostaria de pontuar que é fundamental que o nível de incerteza com relação ao marco legal da IA seja reduzido. Esse é um ambiente que está com muitas incertezas, e essas indefinições são extremamente prejudiciais para os investimentos no nosso setor, no setor de tecnologia do Brasil. Por isso, a gente reconhece a importância de ter essas definições o mais breve possível. Porém, é importante colocar aqui que esse tema é extremamente complexo. É um tema que a gente precisa discutir não só com base ética, mas também com bases técnicas nas nossas proposições aqui.

Quando a gente fala de IA, hoje tudo tem um pouco de IA: quando você abre o seu aplicativo de banco, quando você ouve a sua música no Spotify; manutenção preventiva de aeronaves; até quando você vai ao supermercado e compra um suco de laranja, aquela colheita teve um pouco de algoritmos de IA que contribuíram com aquele suco. Então, a gente consegue ter uma leitura de que, quando a gente está falando de regulamentação da aplicação de IA, inevitavelmente, aqui nesse contexto, a gente está falando um pouco de regulamentação da tecnologia no geral, e não apenas dos usos.

Quando a gente fala de uma forma bem resumida, o escopo da IA é trazer escala e velocidade para as coisas que apenas com o esforço humano eram possíveis até então. Mas hoje a gente está vivendo uma era que é diferente: a IA evoluiu tanto que a gente está na era da criatividade de máquina. Então, é impressionante ver essas novas produções, e é reconhecível também ter certas preocupações acerca de eventuais novos riscos.

Então, o esforço dos Senadores e juristas da Comissão para criar uma proposta justa para o Brasil é evidente – a gente percebe que houve avanços –, mas é importante ressaltar aqui que o texto atual é extremamente preocupante porque ele impacta não só os envolvidos com IA, mas todos os setores produtivos. Como muitos sabem, o PL tem a principal inspiração no modelo da União Europeia, que é reconhecidamente um dos modelos mais restritivos do mundo e é um modelo de lei que está em teste em um mercado que é diferente do Brasil e que possui muito mais incentivos à inovação e ao desenvolvimento tecnológico do que o nosso país. Para dar uma perspectiva, segundo o ITS, que é o Instituto de Tecnologia e Sociedade do Rio, o modelo da Europa, que foi discutido por mais de quatro anos, conta com 39 obrigações. Já o modelo proposto para o Brasil, que a gente está discutindo aqui,



SENADO FEDERAL
Secretaria-Geral da Mesa

contempla pelo menos 69 obrigações, ou seja, a gente pegou o modelo mais complexo do mundo e ainda adicionou 30 novas obrigações.

O trabalho recente da Abria incluiu seis cartas abertas, mais de 20 eventos que a gente fez com a comunidade, e a gente fez um mapeamento de mais de 30 pontos que são preocupantes do PL. Esse foi um trabalho meticuloso realizado em parceria com o time de especialistas do escritório Prado Vidigal. E alguns desses pontos de preocupação... Eu não vou conseguir falar todos aqui, mas, destacando alguns, cito, por exemplo, a explicabilidade. Ela se destaca como um ponto extremamente preocupante porque... Ainda mais hoje, explicar um modelo de IA com um avanço tão grande da inteligência artificial dos novos modelos – a gente está caminhando para o Artificial General Intelligence (AGI) – fica cada vez mais complexo.

Essas tarefas são tecnicamente inviáveis ou extremamente complexas. Isso vai trazer problemas como, potencialmente, expor segredos industriais, ou pior ainda, a gente vai estar expondo para fraudadores. O Brasil é um país que exporta fraudadores infelizmente, e a gente vai estar expondo para fraudadores como funciona um sistema que ele está tentando burlar.

Então, essas altas implicações de governança, Senadores, na prática, vão inviabilizar a operação de pequenas empresas e *startups*, como do nosso colega, como a minha *startup*, como de outros da Abria. Somente os *big players* e as multinacionais vão ter a disponibilidade de tempo, recurso e foco suficientes para manter e cumprir todos os itens de governança que foram mencionados. E aí, a gente vai, inevitavelmente, condicionar o Brasil a uma mera posição de consumidor de tecnologia estrangeira, em vez de ser um desenvolvedor de tecnologia originalmente nacional.

Um outro ponto, como as regras de direito autoral, vai impactar profundamente os desenvolvedores de *startups*, porque vai implicar, do jeito que está redigido, diretamente o treinamento dos modelos, comprometendo o avanço da IA no Brasil. Então, por exemplo, os grandes modelos de linguagem, *large language models*, que estão por trás de ChatGPT, Gemini, que são ferramentas fundamentais, acredito que todo mundo aqui já está usando... Inclusive, isso é estratégico para o Brasil, até para a nossa segurança. A gente precisa de muitos dados, a gente precisa de vídeo, a gente precisa de texto, a gente precisa de foto. Isso para que a máquina entenda o conceito e a relação da linguagem humana, que é extremamente rica. Eu preciso colocar, dando um exemplo aqui, Memórias Póstumas de Brás Cubas, para que ele consiga entender um pouco da nossa riqueza, um pouco do que é a literatura brasileira, e não necessariamente porque eu vou estar replicando aquela obra, ou que eu vá estar fazendo uma cópia daquilo. Quando a gente olha para o Brasil e para o texto, a gente pode também ter algumas inspirações, por exemplo, países como Singapura e Japão, que já previram uma facilidade para treinar esses modelos.

Por fim, Senador, fica muito claro que parece que a gente está mais preocupado em gerar obrigações excessivas, antes dos fomentos necessários para desenvolver o Brasil e, ainda mais, desenvolver o Brasil com uma tecnologia que é apontada como uma das mais promissoras, se não a mais promissora do século, que é a inteligência artificial.

Então, Excelência, este é o momento para a gente parar e pensar para onde a gente quer que o Brasil vá. A gente tem uma proposta muito restritiva, ela pode ter consequências para o futuro do Brasil, não só como uma nação desenvolvedora de tecnologia, mas principalmente na nossa competitividade.

Até fazendo referência também à nossa colega, na agricultura, sistemas de IA já conseguem prever climas, pragas, são responsáveis por um potencial de aumento na produtividade de até 30%. Então, rotular esses sistemas como alto risco – por exemplo, pulverização de agrotóxicos para *drones* –



SENADO FEDERAL
Secretaria-Geral da Mesa

vai comprometer a nossa competitividade e até reduzir benefícios, como a diminuição dos próprios agrotóxicos para o cidadão.

A gente tem que lembrar que o Brasil não é um país fechado, a nossa competitividade não se limita às empresas nacionais. A gente está falando de um mundo onde existem Estados Unidos, China e outros países investindo fortemente em IA, para impulsionar suas economias. Só para você ter uma noção de como a gente percebe que a nossa comunidade está por fora do assunto, tem uma pesquisa de que a Abria fez parte, que é a *Startups & IA Generativa no Brasil*, que acabou de sair, e apontou que somente 14% dos empreendedores estão cientes dos temas sobre os quais a gente está discutindo aqui.

Então, a gente está diante de um tema extremamente complexo, que exige um debate aprofundado, mais participação da sociedade, mais tempo para aperfeiçoar a nossa discussão e, principalmente, ver o que vai acontecer na Europa após esse primeiro ano de implantação.

Sem os ajustes necessários, a gente vai colocar em xeque o futuro de centenas de *startups* que aqui estou representando. E a gente tem uma oportunidade, Senador, de transformar esse projeto de lei em um modelo internacional, em um modelo que vai posicionar o Brasil como líder em IA e acelerar soluções para desafios nacionais que a gente sabe muito bem quais são.

(Soa a campanha.)

O SR. RODRIGO SCOTTI – O incentivo – só para finalizar aqui, Senador – à inovação deveria ser a regra, e não a exceção. O Brasil precisa de empresas competitivas globalmente, o Brasil precisa manter seus talentos aqui e não vê-los indo para empresas gringas. E o Brasil precisa de IAs nacionais, a gente precisa de googles do Brasil, a gente precisa de chatGPTs do Brasil, microsofts do Brasil, a gente precisa de outras empresas desse porte aqui nos representando, e isso só vai acontecer com a união de diversos atores, com a aproximação da academia com o setor privado, com um arcabouço de incentivo governamental para trazer a segurança necessária para o desenvolvimento ético, seguro e promissor da IA no Brasil.

Por fim, Senador, eu gostaria de agradecer imensamente a oportunidade de colaborar com a Comissão. A gente vai disponibilizar todas as referências, o material da Abria, colocar a nossa equipe também para eventuais discussões.

E quero lembrar a todos que a Abria está aqui para colocar o Brasil no mapa global da inteligência artificial...

(Soa a campanha.)

O SR. RODRIGO SCOTTI – ... e esse é o momento de fazer a IA do Brasil acontecer. Obrigado.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Parabéns! Parabéns pela apresentação, Rodrigo!

Sem dúvida nenhuma, em vários pontos interessantes em consonância ou alinhados com o que o Mateus falou, com o que a Patrícia também falou, a gente já começa ver se formar aqui um desenho. É isso que a gente quer, formar esse desenho.

Acho que vocês conseguem entender talvez a minha frustração. Muitas vezes, tendo sido Ministro de Ciência e Tecnologia, tendo estabelecido centros de desenvolvimento de IA no Brasil – oito centros aqui no Brasil –, com uma função que eu vejo... A minha função aqui, justamente, é a de incentivar o desenvolvimento de ciência, tecnologia, *startups*, inovação, transformar conhecimento em nota fiscal, vamos chamar assim, em novos produtos, serviços e empregos no Brasil. E, quando você vê uma



SENADO FEDERAL
Secretaria-Geral da Mesa

legislação que está aqui – eu sou membro, sou Vice-Presidente da Comissão – e você vê que ela está na direção errada, que a gente precisa ajustar, mas tem a pressa de entregar agora, dá para entender a minha frustração talvez de falar isso aqui com vocês. E não tenho nenhum receio de falar esse tipo de coisa porque é assim mesmo. Aqui é uma Casa de discussão, a gente precisa ter vozes de uma posição, da outra posição e discutir a respeito, mas o fato é que os fatos estão levando a ficar muito claro.

Eu espero, sinceramente, que o nosso Relator Eduardo Gomes esteja acompanhando, que acompanhe e que faça as mudanças adequadas para que a gente não perca a oportunidade de colocar o Brasil, como você falou, na vanguarda de uma tecnologia tão importante. A gente tem capacidade intelectual no Brasil para fazer isso, a gente precisa dar o apoio suficiente legislativo para que isso seja possível, sem assustar, sem dar medo, logicamente protegendo o cidadão, o que está previsto na lei, inclusive, naquele substituto que eu falei. Na Emenda 1 ali, você tem todas as defesas previstas para ética, para todo tipo de discriminação, etc., mas sem restringir o desenvolvimento. Essa é a coisa que dá...

E, quando a gente fala de dados, realmente você precisa usá-los para o aprendizado.

Para quem já escreveu ou quem já fez ou já leu, citou artigos científicos, por exemplo, você vai ver que você tem o artigo ali... você, como pesquisador, pode e deve utilizar o trabalho de muitos pesquisadores, livros etc. com a responsabilidade de citar.

Eu estou usando o trabalho do Marcelo. Essa parte veio do trabalho dele. Eu citei. E é isso que tem que ser feito. Agora, você restringir, dizer "você não pode mais usar", isso torna as coisas um tanto difíceis de desenvolver.

E a gente não pode ficar com esse número de restrição dobrado com relação aos países que mais restringem, a gente tem que utilizar o exemplo dos países que mais se desenvolvem, vendo como eles conseguem traduzir isso dentro das leis para atrair pesquisadores, atrair empresas, atrair investimento.

A gente não pode ir para trás. É como querer participar de uma corrida de Fórmula 1 e colocar um motor mil no meu carro para concorrer. Tipo assim, é completamente fora do sentido.

E há insegurança jurídica. Fica muito bom para advogado. Desculpem-me, tem advogado também, mas fica muito bom para advogado. Para quem é engenheiro, você ter alguma coisa que é subjetiva não é adequado de jeito nenhum.

Obviamente, as leis não atendem todo mundo. Sempre alguém não vai gostar de alguma parte da lei. Isso sempre vai acontecer. O que não pode é não gostar de tudo o que está ali. Aí está errado. Aí tem algo muito errado. E a gente precisa ver isso aí.

Eu passo a palavra, agora, ao Jefferson de Oliveira Gomes, Diretor de Tecnologia e Inovação, da Confederação Nacional da Indústria.

Por favor, Jefferson, você tem dez minutos.

O SR. JEFFERSON DE OLIVEIRA GOMES (Para expor.) – Muito boa tarde, Ministro, Senador.

Eu sou Jefferson. Eu sou também professor do Instituto Tecnológico de Aeronáutica, que o senhor conhece muito bem. É a sua casa e sempre será. E coordeno a rede de um dos centros que o senhor criou, a rede lasmim, o centro de Inteligência Artificial Soluções para Manufatura Inteligente, que está dentro do IPT, lá em São Paulo.

Muito obrigado pela sua iniciativa.

Basicamente representando a CNI, fica muito difícil fazer uma apresentação dos modelos que eu tinha preparado, porque nós concordamos exatamente com todas as palavras que o senhor colocou e com a sua brilhante apresentação, Mateus, sobre as restrições específicas a respeito de instalação de novas empresas.



SENADO FEDERAL
Secretaria-Geral da Mesa

Eu vou passar muito rápido.

Também o meu colega Rodrigo muito bem colocou a respeito dessas restrições.

Eu vou colocar, então, para trabalhar esta nossa possibilidade de discussão a respeito do tema, casos que não são só das *startups*, mas basicamente representando as indústrias brasileiras, que eu acho que estão numa vertente.

Eu estou tentando passar aqui a apresentação.

O senhor pode dar um *play* aqui para mim?

Aqui, rapidamente, eu vou passar, porque meus colegas já falaram a respeito disso.

Basicamente, quando o senhor coloca, Senador, a respeito da questão de tempo, é iminente, é determinante essa revolução, e, basicamente, não é agora o momento.

A gente nunca esteve numa situação em que o número de pessoas com mais de 65 anos suplantou o número de pessoas com menos de cinco! Isto é o planeta que a gente está vivendo! É fato essa variação e ainda a menor taxa de natalidade nos últimos 40 anos!

Eu sou de Santa Catarina, no meu estado, a taxa de natalidade está em 0,8%. Dois brasileiros não fazem um novo brasileiro!

Em termos práticos, essa população vai precisar cada vez mais de sistemas com inteligência artificial para auxiliar a vida.

Sobre a nossa indústria, como nós temos pessoas que vão viver mais tempo, cada vez mais, numa única pessoa, você tem quatro, cinco, sete, oito modelos distintos de consumo.

Indústria mais flexível, indústria 4.0.

É impossível indústria flexível sem inteligência artificial. Então, esse é um ponto nevrálgico da nossa conversa para qualquer tipo de questão.

Essa revolução criada por inteligência artificial, para a gente, para o nosso público de 247 mil indústrias brasileiras com mais de cinco funcionários, está restrita a três lugares: claramente a gente não tem condições de trabalhar mais, não somos criadores de infraestrutura para esse processo e vai ser muito difícil nós desenvolvermos *hardwares* e *chips*, mas o que é importante dizer, dentro desse cenário, Senador, é que os sistemas operacionais, os *chips* e as bases de dados ligadas com IA são extremamente mais consumidores de energia.

Em termos práticos, se você vai desenvolver um *datacenter* hoje, com um *chip* especificamente da Nvidia e mais equipamentos específicos dentro desse *datacenter* para gerar trabalhos e inteligência artificial, um único *datacenter*, 3GW de potência. O senhor sabe muito bem que 3GW de potência são mais de 1% da potência declarada que o Brasil tem de energia. Em termos bem simplórios, se o Brasil quer se posicionar também na questão mundial, ele pode ser um lugar para receber *datacenters*.

Como a gente vai trabalhar se formos restritivos à inteligência artificial? É meio ilógico uma situação como essa, sem considerar que também não seremos desenvolvedores de aplicações. O mundo todo está baseado em três, quatro lugares desenvolvedores de aplicações, o que sobra para o resto do planeta são os chamados usuários pioneiros. O agro, a indústria, os serviços, como todos os meus colegas estão falando aqui e utilizam, são usuários pioneiros. Quanto mais cedo você entrar nesse jogo, mais competitivo você fica. É nessa restrição que nós consideramos que seja muito importante.

A gente já mapeou todas as adoções permanentes e estáveis para organizações. Quais são as empresas que aceleram os seus projetos? Em quais países específicos eles estão? Qual a facilidade de uso de cada ferramenta de IA? E sabe qual é o grande problema, no final das contas, nos projetos que nós fizemos junto com o Fórum Econômico Mundial? No final das contas, nós temos três restrições: pessoas mal formadas, pessoas com pouca organização... E a gente ainda está no tempo. Os nossos



SENADO FEDERAL
Secretaria-Geral da Mesa

CEOs estão com idade entre 45 e 55 anos, não são nativos digitais – essa população que vem para frente será nativa – e são muito melhores julgadores de qualquer tipo de problema contra a lei que seja aplicado.

Então, basicamente, para a gente, 80% de todas as ações nas indústrias, Senador, serão modificadas nos próximos cinco anos. E aí a nossa posição é IA para elaboração de projetos, diagnósticos, resumo de implantação de propostas, auxílio à decisão de risco – como o senhor mesmo colocou, a própria IA sendo usada para auxílio à decisão de risco – e filtragem de dados não estruturados; isso basicamente, para a gente, é o que importa.

Vou dar um caso bem simplório para o senhor, Senador.

Em Santa Catarina, nós temos a maior produtora de componentes de engrenagens para motor de arranque do planeta. Se o senhor chegar a essa empresa, no interior de Santa Catarina, o senhor vai encontrar uma indústria com os melhores equipamentos, com as melhores pessoas, organizações e relação na cadeia de valor com todos os países que o senhor possa imaginar, só tem um problema: não teremos mais engrenagens de motor de arranque nos próximos cinco anos. Qual a técnica com essa empresa existente? Como o senhor imagina que nós vamos buscar outras empresas, outros negócios para esse tipo de empresa? Com inteligência artificial, não vai ser mais com projetos de consultoria. Então, filtrar isso, para a gente, baixa basicamente ao impeditivo.

Nós temos o centro de competitividade para produtos brasileiros, que o senhor criou, lá na cidade de São Paulo. Somos uma rede com mais de 89 pesquisadores pelo Brasil afora, e o interessante, Senador, é que todos os projetos tecnológicos, todos esses aqui, que estão sendo colocados nesse quadro, seriam bloqueados, se nós entrássemos nessa lógica – aplicação para empresas.

E o interessante é que essa própria rede que foi criada semana retrasada, num congresso mundial, que aconteceu na França, só sobre segurança cibernética... Porque a gente trabalha desde a lógica de baixo nível, na ponta de robô, que é minha área de trabalho, até a área de lógicas de comunicação entre empresas e segurança cibernética, e o principal trabalho no planeta foi um trabalho do ITA, em primeiro lugar. Brasileiros, em primeiro lugar, desenvolvendo a lógica de como proteger as questões de segurança cibernética.

Perceba que impeditivo bloquearia todas essas etapas que são de mapeamento, plataformas, provas de conceitos, transferências tecnológicas, que é um grande problema que nós temos.

Eu já estou finalizando.

Esse aqui é um futuro da agricultura digital do Brasil.

Nós montamos um centro no Rio Grande do Sul, no Senai, que já trabalhou com centenas de empresas brasileiras que trabalham com o setor de agro, sejam startups desenvolvendo projetos, sejam grandes empresas desenvolvendo produtos.

Curiosamente – não sei se o senhor sabe –, com todas essas universidades lá do Rio Grande do Sul, várias empresas – não sei se o senhor tem a informação –, por mais que o agro seja *pop*, as pessoas não querem trabalhar no campo.

Resultado prático: donos de terras, donos de sistemas agrícolas, aqui no Brasil, estão entrando, porque estão com 70 anos de idade, 80 anos de idade, e empresas brasileiras, aqui no Brasil, como a New Holland, desenvolveram, inclusive, tratores autônomos, para que essas pessoas possam trabalhar.

Tecnologia tupiniquim, com *startups* brasileiras, feita aqui no Brasil – inteligência artificial na veia nesse processo.

Eu finalizo aqui com o B20. Essa questão do B20 é fundamental, porque o Brasil sedia, neste ano, o B20. São os 20 países mais ricos, e o Brasil é responsável pelas escritas relativas ao B20.



SENADO FEDERAL
Secretaria-Geral da Mesa

Um tópico específico é a inteligência artificial.

Nós mapeamos, com todos esses países que fazem parte do B20, aquilo no que cada um está trabalhando, especificamente quais são as preocupações existentes.

Se o senhor olhar naquela figura mais no alto, especificamente, a adoção é o menos problemático; e a gente está trabalhando aqui com questões de restrição à adoção.

Não faz o mínimo sentido que todas essas características relativas a desenvolvimentos nessa área específica sejam bloqueadas.

Eu finalizo aqui a minha apresentação, mostrando um mapa de todos os países...

(Soa a campainha.)

O SR. JEFFERSON DE OLIVEIRA GOMES – ... em qual eles estão. São 37 países, hoje, em desenvolvimento do que nós estamos fazendo aqui, e o Brasil é, de longe – de longe –, o país mais restritivo.

Na última reunião da Mobilização Empresarial pela Inovação, 650 empresários, juntos com a CNI, escreveram um *white paper*, justamente comentando as restrições desse exclusivo PL, a respeito do nosso respectivo trabalho.

Agradeço pela oportunidade e, basicamente, parabéns pela sua condução. *(Palmas.)*

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Jefferson, parabéns pela apresentação.

É bom você ver a sua escola, o ITA, aparecendo também.

Essas coisas que a gente cria como ministro, com uma expectativa de que vão para a frente, quando se fala de desenvolvimento de centros de inteligência artificial, uma coisa tão importante...

Comecei a falar disso aqui lá na transição, em 2018. Quem acompanhou, viu isso aí.

E a gente não pode restringir esse envolvimento; simplesmente não pode restringir esse envolvimento.

Eu estou coletando aqui e, ao final, vou falar um pouco a respeito do processo legislativo também, mas a gente precisa lutar para que isso funcione no nosso país. A gente não pode desistir do nosso país, de forma nenhuma.

E outra coisa, uma coisa que o Mateus falou, durante a sua apresentação, que, por alguma razão, eu me... Não existe crime novo, certo? Aliás, foi uma das coisas que... Em uma das audiências públicas do primeiro lote, vamos chamar assim, a gente estava falando da 2.338 especificamente, e tinha uma jurista que estava exatamente nesse mesmo lugar, ela virou para mim e falou assim: "Senador, o senhor consegue me apontar uma modalidade de crime novo criado pela inteligência artificial?". Eu falei: "Aí você me coloca numa saia justa, eu sou engenheiro... Eu não consigo imaginar". Eu olho de lado, vejo os advogados olhando um para a cara do outro e falo assim: "Não, não tem". Você pode inventar um monte de arma diferente, mas ainda matar é o crime que está previsto no Código. Eu feri alguém, sei lá, essa parte... Você tem essa caracterização.

A gente tem que sair do medo e trabalhar para frente, botar o peito e ter a coragem para ganhar espaço. O Brasil tem capacidade... Você falou de uma coisa também que eu acho sensacional. Nós temos um *grid* de energia no Brasil muito bom, melhor que o da maioria dos países, eu diria que de mais de 80% dos países, e a gente tem o nosso *grid* melhor que o de todo mundo. A gente pode usar essa energia adequada com *data centers*, adequada com inteligência artificial aqui no Brasil. Pensem bem a vantagem competitiva que a gente tem com relação a isso. Aliás, eu tratei disso esta semana com o



SENADO FEDERAL
Secretaria-Geral da Mesa

Secretário de Desenvolvimento lá em São Paulo, também, falando com o Jorge Lima sobre isso, sobre as possibilidades que existem no Brasil como um todo.

Bom, eu passo a palavra agora para a Ana Bialer, membro do Conselho Nacional de Proteção de Dados e Consultora do Movimento Brasil Competitivo, que está conosco remotamente.

Ana, tudo bem contigo? Obrigado por estar aí com a gente, de novo. São dez minutos, e você controla o tempo por aí. Está o.k.?

Obrigado.

A SRA. ANA BIALER (Para expor. *Por videoconferência.*) – Boa tarde, Senador Astronauta Marcos Pontes, boa tarde aos demais.

Antes de mais nada, eu gostaria de me juntar a todos, parabenizando o Mateus pela trajetória dele. Que exemplo para o Brasil! É uma pena Mateus não ter participado das inúmeras audiências que nós tivemos antes, para aprendermos com a história dele. E concordo também com muitas das colocações dele, assim como com a colocação do Rodrigo a respeito da importância do fomento e de a gente olhar para o fomento neste momento. Mas vou tentar me atentar aqui à temática da audiência, no sentido de discutir a classificação de risco.

Antes de mais nada e de avançar, gostaria de parabenizar o Senador Eduardo Gomes pelo trabalho na liderança no Senado, para a construção do marco legal da IA no Brasil e agradecer ao Senador Astronauta por contribuir de maneira tão robusta nas discussões afetas à ciência e tecnologia no Senado, bem como pelo convite para participar deste novo ciclo de debate, representando o Movimento Brasil Competitivo (MBC), para trazer uma visão do setor privado. Eu também sou membro do CNPD, o conselho consultivo da ANPD, mas não falo aqui em nome do CNPD ou da ANPD.

Então, o objetivo da fala hoje é a gente focar na questão da classificação de risco, que está no relatório final do PL 2.338. Em termos de contexto, que eu acho importante a gente trazer, vale lembrar que a abordagem baseada em risco tem norteado todas as legislações afetas à tecnologia nos últimos anos e parece, de fato, ser a abordagem mais moderna. Ela permite balizas normativas a serem adotadas, os *guardrails* a serem desenvolvidos na norma, dando liberdade para o desenvolvimento tecnológico e para a disrupção, permitindo à sociedade colher os frutos do ciclo de inovação, ao mesmo tempo em que protege os indivíduos e a sociedade no geral.

Quando a gente olha para o marco legal da Europa, o EU AI Act é construído com base na premissa de abordagem baseada em risco, estabelecendo uma matriz, fixando riscos excessivos, ou seja, aqueles absolutamente inaceitáveis como sociedade, os de alto risco, os de baixo risco, ou sem risco, ou de risco residual. A estrutura toda do marco europeu é construída para regular esses usos que são definidos como de alto risco, criando um conjunto de regras de transparência, prestação de contas, avaliação de impacto e por aí vai. E, apesar da lista predefinida de alto risco, a Europa tem uma complexa estrutura do que eles chamam de derrogações, deixando claro que há um importante espaço para a criação de exceções, ou seja, ainda que um certo uso se encaixe na lista de alto risco, se a IA for usada sem um risco significativo ou de uma maneira que não seja determinante no processo de decisão, afasta-se essa classificação de alto risco. Eu vou voltar a essa temática quando eu entrar aqui na fala a respeito do que nós temos em termos de Brasil.

A proposta do relatório final do 2.338 no Brasil traz uma estrutura que poderia ser um pouco mais clara, ao trazer para o texto normativo uma pirâmide mesmo de classificação de riscos dos sistemas de IA, a exemplo do que faz o texto europeu. Deixa claro, portanto, quais são as condutas que não podem acontecer de forma alguma, o que nós estamos chamando de risco excessivo – e, como eu disse, essa é uma decisão de política pública propriamente dita –, seguidas daquelas consideradas de alto risco,



SENADO FEDERAL
Secretaria-Geral da Mesa

deixando claro que existe um conjunto de sistemas de IA que possui baixo risco a direitos fundamentais, ou mesmo risco insignificante, ou residual, que estaria, portanto, fora deste arcabouço normativo. Isso parece ser o desejado pela norma, mas não está efetivamente escrito na nossa proposta de texto.

Quando a gente fala especificamente da estrutura que nós temos hoje para risco excessivo, da forma como proposto o texto do art. 13, enquadram-se como risco excessivo o desenvolvimento, a implantação e o uso de sistemas de inteligência artificial em certos casos. Ao tentar ser abrangente, o texto acaba cometendo um preocupante excesso: uma coisa é se proibir que um sistema seja desenvolvido, por exemplo, para fins de disseminar conteúdo de exploração sexual de crianças; outra coisa é se proibir qualquer sistema que possa, eventualmente, ser utilizado para disseminar esse tipo de conteúdo. Não é muito diferente do exemplo já superbatido da faca: a faca não foi desenvolvida para matar; ela pode ser usada por um ser humano mal-intencionado para matar. Então, um ajuste no *caput* do art. 13, para se limitar ao desenvolvimento da IA exclusivamente, seria bastante importante para a gente garantir um ciclo de inovação em termos, pelo menos, de implementação e uso de sistemas de IA no Brasil.

Passando agora ao comentário a respeito do alto risco do art. 14, o texto substitutivo proposto traz uma inovação bastante robusta em relação ao modelo europeu, expandindo na prática, de maneira significativa, os usos que entram como alto risco na legislação brasileira.

Na forma como nós estamos discutindo isso no Brasil, nós temos duas etapas. A primeira é uma lista não exaustiva de usos de IA, o que, na Europa, a gente tem como uma lista exaustiva, efetivamente. Essa lista é o que deverá ser considerado de alto risco, tal como gestão de infraestrutura crítica, alguns usos para educação e formação profissional, lista de espera para serviços públicos essenciais e, até mesmo, administração da Justiça, muito seguindo os moldes do modelo europeu. E a gente cria outra estrutura mais inovadora que traz uma lista de critérios que poderão ser usados, no futuro, pelo SIA, esse sistema regulador de governança que combina tanto a autoridade central quanto as autoridades setoriais, para que eles identifiquem novas hipóteses de alto risco, considerando a probabilidade e a gravidade dos impactos adversos. Essa qualificação de probabilidade e gravidade já foi uma inovação do texto do relatório final que a gente vê com muito bons olhos. Quando se fala em risco, a gente está sempre olhando a probabilidade e o impacto, e isso não era considerado no texto anterior.

O texto do relatório não prevê, de uma maneira estruturada, situações de exceção, aquilo que eu mencionei, há pouco, que a Europa traz como um contexto de derrogações. Embora ele traga a possibilidade de o agente solicitar a exclusão do seu sistema da classificação, ele não traz quais são os critérios que as autoridades poderão utilizar para fazer essa exclusão.

Para que se tenha ideia, o texto europeu prevê, de maneira expressa, a exclusão automática do sistema quando esse sistema tiver o objetivo de realizar procedimentos meramente operacionais; caso o sistema seja usado para melhorar o resultado de uma atividade que tenha sido realizada por um ser humano; terceira hipótese, caso ele seja usado para detectar padrões de decisão, e não para substituir decisões tomadas pelo ser humano; e também quando o sistema for utilizado para atividades preparatórias. Isso chama muito a atenção no item específico de alto risco que a gente tem, que é a administração da Justiça. Quando a gente olha o número de volumes de processos que se tem no Brasil, a adoção de inteligência artificial pode, sem dúvida nenhuma, ajudar e muito a gerenciar esses casos, diminuir o volume e diminuir o tempo dos processos no Poder Judiciário.

Aproveitando que estou falando da administração da Justiça, acho interessante mencionar que, no Fórum de Lisboa, que foi realizado em Portugal, na semana passada, o Presidente do STF, o Ministro Barroso, mencionou uma série de usos que, hoje, o Supremo e a Justiça têm feito, em termos de



SENADO FEDERAL
Secretaria-Geral da Mesa

inteligência artificial, desde sumarizar processos, para que a decisão do juiz possa ser mais rápida, a alguns testes que estão se fazendo, para que se possa traduzir as decisões judiciais de uma maneira mais simples, para que o cidadão consiga de fato entender a complexidade, entender as decisões de uma maneira simples.

O Ministro Barroso concluiu a sua fala, no painel de IA do qual ele participou, com um alerta muito importante que eu acho que beneficia todos nós. Ele falava que a gente tinha que ter uma preocupação, que é a preocupação de não regular em demasia, para não impedir a inovação e não estabelecer uma reserva de mercado para quem já está no mercado. Quando a gente começa a discutir a regulação e estabelecer um patamar muito alto de *compliance*, de obrigações a serem atendidas, a gente acaba, muitas vezes, viabilizando que aqueles já estabelecidos no mercado possam continuar lá, ainda que com um custo mais alto, mas, como já foi falado, a gente dificulta efetivamente a disrupção, o aparecimento de novos unicórnios brasileiros e o benefício da sociedade como um todo.

Então, deixo essa frase do Ministro Barroso para concluir esta fala e trazer, neste momento final de discussão, Senador Astronauta, esses poucos pontos de melhoria, para que a gente possa avançar no que eu entendo que já é uma escolha legislativa, em termos da estrutura de risco, mas que a gente possa fazer esses ajustes pontuais para diminuir potenciais impactos negativos que esse texto traga e ajudar ali no desenvolvimento e na efetiva adoção da inteligência artificial, de modo a ajudar no desenvolvimento do Brasil.

Muito obrigada.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Obrigado, Ana.

Ana Bialer é membro do Conselho Nacional de Proteção de Dados e Consultora do Movimento Brasil Competitivo.

Entre vários pontos importantes... Para quem trabalha com risco, como eu falei no início aqui, quando a gente fala de um cenário em que você tem uma possibilidade – quero lembrar também que a gente sempre usa risco no sentido negativo, mas também pode ser uma oportunidade, o pessoal se esquece dessa outra parte –, o risco sempre está relacionado, de alguma forma, com a probabilidade de algum evento acontecer que seja, vamos supor aqui, negativo e também, se aquela probabilidade se confirmar, com qual o impacto daquilo. Mas aí vem a próxima pergunta: o impacto para quê? É uma pergunta que a gente sempre tem que fazer. A gente fala aqui de ética, discriminação, uma série de possibilidades negativas da utilização, e, dentro de um contexto de análise de incidentes ou acidentes, você sempre trabalha com o que a gente chama de fatores contribuintes. E aí, dentro desses fatores contribuintes, vamos dizer assim, é onde moram os demônios. Como a utilização de algo pode causar um fator contribuinte ou uma sequência de fatores contribuintes que vai levar a um incidente ou um acidente? Quando a gente pensa em risco, é bom ter isso aí em mente, porque, se você tentar fazer qualquer tipo de regulação que vai prevenir absolutamente tudo, muita coisa vai ser ilógica, injusta, distorcida, e isso pode causar efeitos completamente inesperados e negativos dentro do sistema. Então, estou falando com uma linguagem um pouco de engenheiro aqui, mas é só para se ter ideia de como isso é complicado se você não prestar atenção em que tipo de exigências você coloca dentro de uma legislação como essa daí ou que tipo de decisão vai ficar na mão de quais pessoas, com quais qualificações, para realmente decidir sobre aquilo lá. Só para se ter isso em mente.

Só confirmando aqui... (Pausa.)



SENADO FEDERAL
Secretaria-Geral da Mesa

O.k., eu passo a palavra então, na sequência, ao Marcelo Almeida, Diretor de Relações Institucionais e Governamentais da Associação Brasileira de Empresas de Software, que congrega um número muito grande de empresas, desde pequenas até muito grandes.

O SR. MARCELO ALMEIDA (Para expor.) – Exmo. Sr. Senador Astronauta Marcos Pontes, muito obrigado, mais uma vez, pela oportunidade.

Cumprimento aqui o Mateus, que está ao meu lado, e, na pessoa dele, cumprimento todos os meus colegas que estão aqui.

Eu tinha preparado, Senador – não é a primeira vez que eu trato deste projeto aqui no Senado Federal –, um roteiro, mas eu acabei ficando depois das falas dos meus colegas e, assim como a inteligência artificial vai se alimentando de dados, de *inputs* de dados que vão sendo recebidos, eu vou me permitir fazer algumas adaptações a partir dos dados que eu fui recebendo da fala dos meus colegas aqui.

Parece-me que está muito claro que a gente tem uma diferença paradigmática de visões que estão espelhadas neste Projeto de Lei 2.338 sobre aquilo que nós entendemos, portanto, que deve ser uma regulamentação de inteligência artificial no Brasil. E todas as falas estão convergindo no sentido de considerar que aquilo que está sendo chamado de riscos excessivos, na verdade, é um excesso – e esse verdadeiramente é o risco –; aquilo que está sendo chamado de avaliação preliminar, na verdade, é uma censura; e aquilo que está sendo chamado de alto risco, na verdade, é aquilo que nós deveríamos assumir como o risco de ter uma inteligência artificial brasileira competitiva, inovadora e que traga benefícios às mais variadas ações econômicas – a exemplo do que a gente pode ver, como trouxe muito bem aqui o Jefferson, na área da agricultura.

Eu vou pegar essa sua fala, Jefferson, para dizer que talvez o que nos falte aqui seja, efetivamente, trazer mensagens como esta que você trouxe sobre a importância de regulamentar isso quanto aos efeitos – V. Exa. muito bem falou no início – que vão atingir invariavelmente a cadeia produtiva nacional. Qualquer que seja o setor vai ser atingido pela regulamentação de inteligência artificial, seja a agricultura, o setor financeiro, a indústria, os produtores de conteúdo... Todos eles vão ser atingidos pela regulamentação de inteligência artificial. Portanto, é importante nós chamarmos a atenção desses agentes que são determinantes na economia, para que possamos convocá-los a prestar atenção no debate e, definitivamente, chamá-los a mobilizar os seus respectivos representantes para trazerem ações para isso.

Dito isto, Senador, eu quero me propor aqui a analisar três projetos como paradigmas da minha análise. O primeiro é o Projeto 21. O segundo é o Projeto 2.338 – para o qual, salvo engano, é o terceiro ou o quarto relatório que o Senador Eduardo Gomes apresenta, o que, por si só, já revela a necessidade, talvez, de tempo para o amadurecimento da matéria. O próprio Senador Eduardo, ao longo da sua relatoria, vem trazendo mecanismos de, digamos assim, evolução do texto, mas digo novamente: acho que nós temos um problema nesse aspecto, porque eu tenho discordâncias paradigmáticas com isso. E o substitutivo de V. Exa., que é a Emenda 1, uma emenda substitutiva com relação ao texto.

A primeira coisa que eu gostaria de dizer é, mais uma vez – como eu disse nos debates temáticos –, sobre o tempo. Se quiséssemos, então, aprovar uma lei mais rapidamente, em que o tempo efetivamente fosse o verdadeiro motivador para que tivéssemos uma lei regulamentadora da inteligência artificial, talvez, sob a ótica legislativa, devêssemos partir do 21, porque, afinal de contas, é um projeto que já foi aprovado pela Câmara dos Deputados, revisado, então, numa perspectiva inicial, pelo Senado, que o encaminharia, portanto, seguindo o rito constitucional de projeto de lei ordinário, para a sanção da Presidência da República.



SENADO FEDERAL
Secretaria-Geral da Mesa

Então, se o tempo for o paradigma para que a gente possa refletir, talvez tenhamos que considerar que, no contexto legislativo, a gente já tem um projeto aprovado na Câmara, basta revisá-lo no Senado e mandá-lo para a Presidência da República, mas não me parece que esse seja o cenário posto. O cenário posto é um convite à criação de direitos, deveres e obrigações, muitas vezes, com desprestígio de inovação e desenvolvimento tecnológico.

E aí, nesse exercício, a gente tem, de fato, uma interferência – parece-me – verticalizada, muito assertiva para a criação de direitos, poucas inovações e poucos incentivos para que a gente possa verdadeiramente desenvolver a tecnologia de inteligência artificial no Brasil.

A primeira divergência está até no conceito, Senador, de sistemas de inteligência artificial. Veja que o 21 fala em "um conjunto de objetivos definidos por humanos [que] pode, por meio do processamento de dados e de informações, aprender a perceber e a interpretar o ambiente externo [...]". Então, o verbo aqui é "aprender".

O 2.338 diz que o sistema de inteligência artificial é um sistema baseado em máquina com graus diferentes de autonomia para objetivos explícitos ou implícitos inferir, a partir de um conjunto de dados e informações, o que recebe para gerar resultados, em especial, previsões, conteúdos, recomendações ou decisões que possam influenciar o ambiente virtual, físico ou real. Aqui, o verbo é "inferir".

E o substitutivo de V. Exa. diz assim:

Art. 4º Para as finalidades desta Lei, adotam-se as seguintes definições:

I – sistema de inteligência artificial: sistema computacional, com graus diferentes de autonomia, desenhado para inferir como atingir um dado conjunto de objetivos, utilizando abordagens baseadas em aprendizagem de máquina [...] lógica [...] [de] representação do conhecimento, por meio de dados de entrada provenientes de máquinas ou humanos, com o objetivo de produzir previsões, recomendações ou decisões que possam influenciar o ambiente virtual ou real;

.....
V. Exa. teve a sabedoria de agregar as diferentes previsões sugestivas do legislador para poder tornar o conceito um pouco mais abrangente, contemplando, portanto, esses dois cenários.

Então, quando nós tratamos de sistema de inteligência artificial, parece-me que inferência é um dado importante, que pode ser, por sua vez... eu estou aqui com o atento olhar do Marcelo Moraes, que, cientista como é, é quem melhor classifica as questões de inferência. Um cientista tem que saber fazer inferências, sejam elas dedutivas, sejam elas indutivas, para que a gente possa ter a inteligência artificial, portanto, gerando todos os seus benefícios. A dedutiva parte de fatos estabelecidos e regras para poder fazer com que a inteligência artificial gere os seus efeitos. A indutiva, por sua vez, aprende os dados para fazer previsões ou tirar conclusões.

É lógico, Marcelo, com a licença que eu peço a você, eu estou simplificando um conceito que é absolutamente complexo, mas que, dentro desse cenário de conceituação de sistemas de inteligência artificial, partindo, portanto, desse mecanismo de inferência indutiva ou dedutiva, nós temos, sim, um ambiente para que, a partir de uma coleta de dados, a partir de um pré-processamento, a gente possa, portanto, construir um modelo de inteligência artificial que gere resultados, gere essas inferências, que vão depois passar por um processo de avaliação e refinamento, muitos deles executados pelos humanos propriamente, que é quem toma as decisões.

Isso é absolutamente importante de a gente dizer, porque muitas vezes, a gente olha o Projeto 2.338 e a gente vê uma série de preocupações, de vieses, uma série de preocupações, de ações que, muitas vezes, são protagonizadas por seres humanos, e não necessariamente pelas máquinas. Então a



SENADO FEDERAL
Secretaria-Geral da Mesa

gente precisa observar a realidade como ela é, e não tentar espelhar obrigações, direitos e deveres com realidades que a gente gostaria que fosse. Mas eu também gostaria que a nossa realidade fosse mais igual, mais digital e menos desigual, que é o lema da Abes, que eu represento aqui. Mas isso é um processo que, muitas vezes, não se faz por mecanismos legislativos verticalizados e impostos.

(Soa a campanha.)

O SR. MARCELO ALMEIDA – Faz-se por uma construção legislativa que tem aqui um *start* daquilo que eu chamo de governança legislativa. Já falei isso em outros fóruns, Senador, e fui muito criticado, porque parecia que eu estou promovendo uma legislação que não regula nada, que só prevê as circunstâncias e não funciona para nada.

Só que o processo legislativo não acaba nessa etapa. Ele é um processo que pode gerar os *insights* necessários para que a gente possa ter mecanismos reguladores mais próximos da realidade avizinhada da tecnologia posta, e aí, sim, a gente trazer efetividade para a regra normativa vigente.

Como V. Exa. já também disse aqui, e é muito importante a gente sempre pontuar, a gente não pode criar marcos legislativos com a rigidez processual constitucionalmente estabelecida para que a gente possa ser provocado a ficar revisando-a ano a ano. Projeto de lei não funciona para isso.

(Soa a campanha.)

O SR. MARCELO ALMEIDA – Projeto de lei funciona como um elemento que prevê norteadores fundamentais para que todo o sistema que eu chamo de governança legislativa, por intermédio de decretos e tudo mais, possa estabelecer aquilo que efetivamente é arriscado na sociedade brasileira e que precisa ser protegido. E não presumir os riscos, pré-qualificando todos eles nos arts. 12 e 13, para dar um exemplo concreto, falando de riscos excessivos ou de alto risco, para que a gente possa ter, aí sim, um cenário de qualificação em que, dentro desse ambiente de governança legislativa, a gente possa trazer o resultado.

E, para terminar, Senador, V. Exa. pediu para a gente fazer propostas. Então eu acho que nós temos na mesa hoje, no contexto legislativo, dois projetos, como disse, como paradigmas dessa minha fala: o substitutivo de V. Exa. e o Projeto 21, a partir do qual deve nascer um novo substitutivo...

(Soa a campanha.)

O SR. MARCELO ALMEIDA – ... considerado todo esse contexto de discussões que está posto aí, para que nós possamos sugerir para a sociedade brasileira mecanismos que sejam assertivos, dentro da lógica que é posta para o desenvolvimento de inteligência artificial transnacional, diga-se de passagem.

A inteligência artificial não obedece a fronteiras. Então a gente precisa considerar a realidade mundial, mas trazer para dentro do contexto brasileiro um ambiente de absoluta tranquilidade ao desenvolvimento tecnológico, partindo desses dois projetos, na minha sugestão a V. Exa., para que a gente possa oferecer substitutivos e, aí sim, contribuindo para a governança legislativa, a gente poder atingir resultados eficientes.

Muito obrigado.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Obrigado, Marcelo.

Você vê que as visões são muito parecidas aqui, no sentido do que é importante.



SENADO FEDERAL
Secretaria-Geral da Mesa

Aliás, enquanto você falava, eu estava raciocinando aqui com uma coisa: tem muita gente que não sabe das capacidades da inteligência artificial, e, portanto, perde em eficiência, perde em competitividade e assim por diante.

Por outro lado, tem muita gente que imagina que a inteligência artificial tenha capacidades que ela não tem. Acho que o pessoal imagina que vai aparecer um Schwarzenegger aqui, vestido de exterminador do futuro. Ela não tem esse ponto, ela não tem essa capacidade ainda.

O fato é que o desconhecimento gera medo, o medo gera decisões ruins – e decisões, geralmente, restritivas. É o que a gente tem que evitar, a gente tem que trabalhar com fatos.

Então, eu concordo que a gente precisa ser pragmático para encarar essa situação que a gente tem pela frente. E, de novo: vocês já sabem muito bem a minha posição. Vai continuar sendo dessa forma. Infelizmente, não sou eu que tomo todas as decisões, não tomo as decisões sozinho, isso faz parte de um colegiado.

Eu espero... Depois vou fazer uns comentários finais, mas eu espero que os nossos Senadores – que estão nos acompanhando – acompanhem; ou que as assessorias passem para os Senadores essas informações, para que eles usem a cabeça, o bom senso e a nossa responsabilidade de trabalhar a favor do Brasil, não contra o Brasil. Trabalhar contra o desenvolvimento de tecnologia no país é trabalhar contra o Brasil. Então, a gente não pode fazer isso aqui, de jeito nenhum.

E há outros riscos também, que são... Como se falou, o nosso medo é o maior risco agora. Isso é impressionante. Lembro até uma frase de um Presidente americano, lá atrás: "O maior medo é o medo em si" – algo desse tipo. A gente não pode deixar o medo nos vencer. A gente tem que vencer isso aí e trabalhar da forma correta.

Eu passo a palavra agora, por dez minutos, ao Felipe França, Diretor-Executivo do Conselho Digital.

Por favor, Felipe. Dez minutos.

O SR. FELIPE FRANÇA – Obrigado, Senador, pela oportunidade.

Obrigado a todos os meus colegas que contribuíram. É aquele velho clichê: é sempre difícil falar depois da oratória do Marcelo; dos casos do Mateus e do Rodrigo, que conseguiram transformar uma coisa concreta, a vida do *startup* brasileiro, que está inovando, se desenvolvendo com inteligência artificial; da experiência do Jefferson nas aplicações, que são tão diversas.

Quando a gente pensa em inteligência artificial, muitas vezes, a gente pensa só na tecnologia, mas a gente está tendo uma discussão que é absolutamente transversal, que impacta toda a economia e a vida de todos os indivíduos. A gente está com inteligência artificial em absolutamente tudo da nossa vida hoje em dia: o *spam* do nosso *e-mail*, a coleta na hora do suco... Está em tudo.

E há também às colegas, a advogada do Peck – foi sensacional a apresentação; e a nossa amiga Ana Bialer, que também fez uma apresentação maravilhosa.

Eu acho que eu não vou trazer nada absolutamente novo, mas vou reforçar alguns pontos que os amigos falaram. E o primeiro deles é que é importante regular os riscos, e não a tecnologia.

Por mais que a gente tenha discutido um projeto que se diz – analisando – um modelo baseado em risco, na realidade, quando a gente vai ler um pouco o texto, e eu faço questão de trazer o texto para a gente parar para analisar, a gente está regulando, na verdade, a tecnologia.

Se a gente ler o art. 14, por exemplo, quando ele fala: "Consideram-se sistemas de inteligência artificial de alto risco aqueles desenvolvidos e utilizados para as seguintes finalidades e contextos:" E eu vou dar um exemplo aqui: "XI - sistemas de identificação e autenticação biométrica para o



SENADO FEDERAL
Secretaria-Geral da Mesa

reconhecimento de emoções, excluindo-se os sistemas de autenticação biométrica cujo único objetivo seja a confirmação de uma pessoa singular específica".

Eu acredito que, quando o Relator fez o projeto, estava mirando aqui: como é que se vai fazer o reconhecimento facial, quais são os impactos e tudo mais. Mas, na realidade, ele também pode estar acertando aqui. É uma tecnologia que, da mesma forma, está fazendo a identificação e a autenticação biométrica para reconhecimento de emoções. Isso aqui é basicamente um sistema de inteligência artificial, um filtro numa rede social.

Voltando atrás para me apresentar, eu sou Felipe França, sou Diretor-Executivo do Conselho Digital. O Conselho Digital é uma entidade sem fins lucrativos que representa, estuda e coordena o ecossistema de aplicações de internet.

Esse é um dos exemplos e a gente poderia falar de como vai afetar a vida nessas aplicações de internet. Mas a gente também poderia falar de outro: *os feeds*, por exemplo, de uma rede social que você está utilizando agora, se qualquer pessoa pegar no seu celular. Eu não sei exatamente em qual ponto que o Relator estava mirando quando ele colocou, no inciso XIII:

XIII – Produção, curadoria, difusão, recomendação e distribuição, em grande escala e significativamente automatizada, de conteúdo por provedores de aplicação, com o objetivo de maximização do tempo de uso, engajamento das pessoas e grupos afetados.

Não sei exatamente qual é o risco que ele está querendo. Pode ser desinformação e aí, se a gente der um Ctrl+F no relatório, a gente vai encontrar algumas vezes a questão da integridade da informação. Talvez seja esse o interesse, talvez seja discutir questão de discurso de ódio. Não sei exatamente qual o objetivo. Mas a gente sabe que quando ele está colocando essa tecnologia da produção, sistemas de recomendação, distribuição, também está mirando no *feed* das plataformas. Há algumas décadas, quando você está rolando o seu *feed*, ele não é mais cronológico, ele está vendo o que é mais sua cara. Eu gosto de futebol, gosto de viajar, vão aparecer para mim futebol e viajar, talvez não apareça para mim culinária.

Então, é importante a gente entender: a gente está regulando aqui a tecnologia e vai ter efeitos na nossa vida. Por ser um sistema de alto risco, vai estar sob o SIA, sob a ANPD, e vai ser regulado lá esse sistema. Então, esse é um primeiro ponto em que a gente precisa prestar atenção: a gente está regulando a tecnologia e não o risco. Isso vai afetar nosso dia a dia, inclusive em questão de liberdade de expressão.

Eu entendo... O Senador Eduardo Gomes falou, várias vezes, que não é interesse dele tratar, por exemplo, de questão de desinformação, que é um tema que está sendo discutido no PL 2.630, lá na Câmara. Mas, de certa maneira, isso está sendo trazido ao debate quando coloca um inciso como esse.

Outra questão que a gente queria falar – e aí eu vou tentar não ser longo no debate – é custo e conformidade. Vou tentar discutir um exemplo que foi trazido por alguns colegas, que é o exemplo da explicabilidade. Separei alguns trechinhos. No art. 3º ele fala:

Art. 3º O desenvolvimento, a implementação e o uso de sistemas de inteligência artificial observarão a boa-fé e os seguintes princípios:

.....

VI – transparência, explicabilidade [...] [observado o segredo comercial e industrial];

Legal. No art. 6º, ele também vai falar lá de direito à explicação para decisão e, no parágrafo único, vai citar várias granularidades de como é que vai funcionar essa explicabilidade.



SENADO FEDERAL
Secretaria-Geral da Mesa

Em tese, é uma coisa muito interessante, a gente quer que tenha explicabilidade, quer que seja uma coisa o mais transparente possível, mas direitos implicam escolhas. Existem *trade-offs*. E aí, conversando com o Senador, que é engenheiro, ele está entendendo do que a gente está falando aqui: eu posso ter um algoritmo mais efetivo e menos explicável e, do outro lado, um algoritmo menos efetivo e mais explicável, a depender da complexidade. Se eu for tratar, por exemplo, de um ChatGPT, de um modelo de linguagem, ele vai ter milhões de parâmetros e vai se tornar menos explicável. E aí, se a gente pensar em questões mais complexas, por exemplo o clima... O clima não é simples de explicar.

Você vai pensar num sistema de inteligência artificial que ajude a gente a identificar quando vai ter um alagamento. Ele vai envolver um modelo hidrológico, um modelo de inundação, a distribuição. É complexo. É muito complexo. E aí a gente vai querer um sistema mais explicável ou mais efetivo. E eu não estou dizendo que um elimina o outro, mas a gente precisa ter noção de que existe esse *trade-off*; quanto mais complexo, menos explicável na linguagem natural. E isso é uma decisão.

Eu estou falando isso aqui só como um exemplo de custo de conformidade. A gente poderia falar de várias outras questões, que, de alguma forma, vão afetar a forma como os desenvolvedores brasileiros vão trabalhar com isso.

O terceiro ponto que eu queria tratar é sobre disponibilidade de dados. E eu sei que esse é um ponto muito delicado. Ele é muito delicado porque a gente precisa achar um ponto de equilíbrio. A gente precisa entender que a gente precisa proteger o titular do direito na criação dele. Ao mesmo tempo, a gente precisa entender que o Brasil está querendo inovar, se desenvolver e ser competitivo em termos globais. E aí é bom a gente ter um pouco da história.

Como é que o GPT surgiu? Quando o primeiro ChatGPT surgiu, 80% do *corpus*, como a gente chama, o que estava disponível de dados para ele analisar, veio do Common Crawl, que é uma associação americana que fez a raspagem de mais de 250 bilhões de páginas durante 17 anos na internet. E isso não é medido em *megabyte*, não é em *gigabyte*, não é em *terabyte*; é em *petabyte*. Cada um é 1024 alguma coisa de multiplicação. Os senhores engenheiros me ajudem, depois, com essa resposta. Mas, assim, é uma quantidade imensa. E é muito difícil, quando você está tratando de *petabyte* – e o GPT surgiu porque ele conseguiu ter acesso a todos esses dados para conseguir se desenvolver –, você conseguir explicar isso.

Então, nos Estados Unidos, por exemplo, eles têm, jurisprudencialmente, a regra do *fair use*, que aqui também é válida hoje. No Japão, em Singapura, em Israel, eles desenvolveram medidas legislativas para garantir uma segurança jurídica a quem estava inovando, sobre quando é e quando não é uma violação ao direito do autor. E isso não é uma questão simples, porque a gente precisa proteger o autor. Ao mesmo tempo, precisa garantir disponibilidade de dados. Quando não feriria? E aí, no texto fala "direito de autor e conexos". Aí tem todo um capítulo, acho que do 60 ao 64 no relatório atual. Vai-se modificar, então não sei exatamente como é que vai ficar, mas ele fala: "o desenvolvedor de inteligência artificial que utilizar conteúdo protegido por direito de autor e conexos no seu desenvolvimento deverá informar quais conteúdos protegidos foram utilizados no processo de treinamento do sistema de inteligência artificial, conforme disposto em regulamentação". Então, ele pede que você diga quem é um, proíbe para que sejam fins comerciais. Eu fico imaginando como se faria isso com Common Crawl, que são 90 *petabytes*, 250 bilhões. E a resposta não é simples.

Eu sei que nesse momento deve ter um monte de gente no *chat*, agora, batendo em mim, porque eu estou discutindo esse assunto. A resposta, de fato, não é simples, mas ela foi tratada – eu vou passar aqui depois –, ela foi tratada no PL 21 que foi aprovado na Câmara. No art. 5º, ele colocava: "São princípios para o desenvolvimento e aplicação de inteligência artificial no Brasil..."



SENADO FEDERAL
Secretaria-Geral da Mesa

(Soa a campanha.)

O SR. FELIPE FRANÇA – ... disponibilidade de dados: não violação de direito do autor pelo uso de dados, de bancos de dados e de textos por ele protegidos, para fim de treinamento de sistemas de inteligência artificial, desde que não seja impactada a exploração normal da obra para o seu titular". É uma solução relativamente simples. Se você não está afetando o outro lado, você conseguiria resolver isso.

Então, se você utiliza um episódio de Friends e Sex and the City para aprender a falar inglês novaiorquino, até aí não tem... Mas, se você utiliza isso para fazer um roteiro de uma série, sim, seria um problema. Então, como fazer que não tivesse esse choque.

E, aí, isso é muito importante para a gente ter acesso a textos em português. Se a gente não tem dados em português, você não conhece a realidade brasileira. A gente sabe que, hoje, tem sistemas de inteligência artificial, modelos de linguagem, rodando na Europa sem dados europeus. Eles não confiam na segurança jurídica de poder utilizar dados europeus. Eles não vão saber a regionalidade. E eu digo isso como pernambucano que mora em São Paulo, com uma filha carioca...

(Soa a campanha.)

O SR. FELIPE FRANÇA – ... a gente tem particularidades regionais gigantescas que a gente precisa entender. Quando é que ele vai saber que um paraense, quando está falando "égua", isso é uma interjeição? Ele precisa ter acesso a esses dados. A gente não vai saber a realidade, inclusive para você saber quando, naquele contexto, o que o brasileiro está fazendo. Então, a gente precisa de dados brasileiros para entender a realidade brasileira.

Queria encerrar com uma última pergunta: qual o objetivo do PL nº 2.338, de 2023? Eu acho que essa é a pergunta-chave. O objetivo é fomentar, fazer o Brasil ser campeão na batalha da inovação? Se for, eu acho que os meus amigos, o Scotti e o Mateus, responderam bem. Ele não está conseguindo cumprir esse papel. O MID chegou a informar – eu acho que no relatório anterior – quantos por cento tratavam de fomento à inovação. Era coisa de 5%. A gente não conseguiu.

Se o objetivo é trazer segurança para o usuário, não sei se ele está conseguindo trazer essa tranquilidade também. Se é combater a *deep fake*, isso está na resolução do TSE. Está proibido já! Não tem para que ter esse medo.

Não está claro, para mim, qual o objetivo do 2.338 como ele está escrito hoje, mas, definitivamente, não está trazendo segurança jurídica para quem está querendo inovar, não está trazendo segurança jurídica para quem está no estado, o estado-juiz, e por aí vai. Então, a gente precisa responder essa pergunta.

Eu agradeço muito ao Senador por dar esta oportunidade para a gente. Espero que os demais Senadores também possam ver.

Obrigado.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Obrigado, Felipe.

Sem dúvida nenhuma, é isso aí – sem dúvida nenhuma! Aliás, vou começar pela parte de trás, falando sobre essa questão do risco. Como eu já falei outras vezes, é muito mais arriscado não usar ou restringir demais a inteligência artificial do que usar, mesmo porque, para qualquer um que já trabalhou com mitigação de risco, se você não conhecer efetivamente o risco, categorizar e conhecer detalhes daquele risco, você nunca vai ser capaz de fazer uma mitigação descente daquele risco. E a pior maneira



SENADO FEDERAL
Secretaria-Geral da Mesa

de você fazer é se afastar do desenvolvimento, seja lá daquele sistema, ou deixar por conta de outros. E, aí, como é que você vai mitigar um risco sobre o qual você não tem controle? É bom ter isso aí em mente, sem dúvida nenhuma.

Outra coisa, muito bem colocada aí para quem conhece desenvolvimento de aprendizado, é você fazer essa relação de que eu vou ter que explicar como que eu cheguei a uma certa conclusão. Colocando em termos simples aqui – vou colocar assim para o pessoal que está assistindo acompanhar também: através de uma quantidade enorme de dados, a inteligência artificial chegou, através do treinamento, a uma solução daquilo lá. Esse sistema funciona dessa forma. Explicou. Agora, você vai explicar como ela chegou ali.

Aliás, para quem pensa muito em usar inteligência artificial... Às vezes, a gente pensa geralmente só na parte complexa da inteligência artificial, mas, se for pensar em sistemas mais simples, ela já começa lá com regressão, regressão linear, que a gente faz o tempo todo. Quem trabalha com matemática, física, engenharia ou qualquer outra coisa faz regressão linear. Você tem dados rotulados, você mede 10x aqui, vai dar 10y lá, você vai ter vários pontos espalhados, passa uma reta com mínimos quadrados – vamos chamar assim, para a diferença quadrada, para não dar mais e menos –, passa uma reta ali no meio e, a partir dali, você pode ter um X, que você não mediu, mas você tem uma previsão de quanto vai ser o Y, baseado nos resultados que já teve antes. Dados rotulados, regressão linear... ninguém chamava isso de inteligência artificial na época, mas é um dos roteiros – vamos chamar assim – dos algoritmos de inteligência artificial. Esse você consegue explicar, esse é fácil de explicar.

Aí vai lá para *deep learning*, pega uma rede neural convolutiva, por exemplo, análise de imagem. Você vai começar a buscar com várias camadas, copia mais ou menos o que nossas sequências de neurônios fazem no córtex, e vai passando informação de uma para outra, simplificando; entra uma retificação linear para, depois virar vetor, aquela coisa toda. Vai explicar isso aí! Eu consigo explicar o processo. Agora, como isso vai acontecer com os dados concretos ali, na hora, não dá para fazer isso aí; ou seja, quem advoga uma situação dessas é basicamente porque não conhece o sistema; não tem conhecimento. Aliás, esse é um problema que a gente tem em muitas coisas: são pessoas tomando decisões sobre coisas que não conhecem. Isso é fatal para o desenvolvimento de qualquer coisa.

Outra coisa é a integração de informações. Nessa questão de integridade – melhor, de informações –, a gente começa a entrar numa seara que é *fake news*, etc. É um problema sério, mas tem legislação própria para tratar disso; isso não é para ser tratado pela legislação de inteligência artificial. Tem outras coisas para fazer isso aí. Assim como a gente não vai tratar de proteção de dados; tem a Lei de Proteção de Dados também, a gente tem que se alinhar com ela. Então, você vê que vários pontos acontecem nesse sistema.

Agora, nós passamos por todos os nossos oradores. Eu queria ler... vou fazer a leitura aqui das perguntas e dos comentários que nós recebemos através do e-Cidadania. Depois, vou devolver a palavra novamente para cada um dos oradores.

Eu vou pedir então, quando devolver a palavra, para a gente fazer, dado o adiantado da hora, ao mesmo tempo as considerações finais. Eu vou dar três minutos para cada um fazer as considerações finais e responder a alguma, ou uma, ou duas, algumas das questões ou comentários que estão aí. Então, já vão pensando em como encaixar isso nos seus comentários finais.

Eu vou começar pela mesma ordem que nós fizemos, com exceção daqueles que tiveram que sair, como o Mateus, que tinha uma viagem – ele teve que sair antecipadamente –, assim como a Patricia. Eu vou passar a palavra novamente na sequência que nós fizemos. Então, vou começar com o Rodrigo.



SENADO FEDERAL
Secretaria-Geral da Mesa

Rodrigo Scotti, Presidente do Conselho de Administração da Associação Brasileira de Inteligência Artificial, para responder algumas das perguntas que quiser e fazer suas considerações finais em três minutos. É um desafio, eu te digo, mas é importante por causa do nosso tempo.

O SR. RODRIGO SCOTTI (Para expor.) – Bem, eu gostaria de reiterar a validade das preocupações de todos, principalmente com relação a questões de proteção das suas obras, dos direitos autorais, de artistas, de redatores e pessoas que estão preocupadas sobre como vai ser o futuro dos trabalhos, principalmente acerca de uma inteligência artificial gringa conseguir se inspirar e roubar esses trabalhos.

Justamente por isso, a gente colocou aqui, em vários pontos, que é uma questão complexa. Todos os pontos colocados aqui trazem apontamentos sobre o PL, mas é complexo sair com soluções prontas, porque isso está sendo prototipado no mundo exatamente agora.

Então, a gente não quer, por exemplo, que um modelo de linguagem aprenda Friends e faça uma cópia do roteiro de Friends, mudando apenas alguns aspectos; a gente não quer que a inteligência artificial vá roubar fotos de artistas e vá gerar imagens que são muito próximas daquilo. São extremamente compreensíveis todas essas preocupações.

Eu acabei de receber novas perguntas aqui e eu não consegui ler tudo em detalhe, mas eu acho que esses pontos que vão na minha consideração final são: primeiro, é importante a gente entender que – assim como o senhor falou, Senador – muito do que se trata com relação à inteligência artificial está sendo discutido por pessoas que têm pouco conhecimento técnico de como funciona a inteligência artificial. Isso é um ponto extremamente relevante, porque hoje em dia você abre o seu LinkedIn e todo mundo é especialista em AI, todo mundo é especialista em inteligência artificial; está todo mundo querendo entrar nesse barco e se aproveitar um pouco dessa onda...

(Soa a campainha.)

O SR. RODRIGO SCOTTI – ... colocando-se como especialista, e, na verdade, é uma questão técnica que a gente tem que entender – o senhor mesmo colocou aqui alguns aspectos – que, ora, a gente está no começo de uma revolução, uma nova revolução, e a gente tem que ter essa consciência. Os artistas, proprietários de direitos autorais, têm que também entender que a gente está vivendo um novo mundo; o mundo não vai voltar para trás. O Brasil não vai virar um país fechado, a gente não vai se fechar.

A gente tem que ter uma legislação que é compatível com o desenvolvimento de inteligência artificial mundo afora, e a gente precisa entender, então, como a gente vai abarcar todas essas preocupações, e não vai ser apressadamente – e é esse o ponto dos colegas: ninguém quer, aqui, fazer grandes proteções de *big players*; pelo contrário, até coloquei aqui o meu ponto de que, com a regulação do jeito que está, somente quem vai sair ganhando com isso serão os *big players*. E não é que saiam ganhando; na verdade, vão sair perdendo menos...

(Soa a campainha.)

O SR. RODRIGO SCOTTI – ... porque eles já têm estrutura para isso.

Então, a gente precisa – só fechando aqui, Senador, desculpe-me – olhar para referências internacionais, olhar para o que está sendo feito, respeitar também os artistas, respeitar os autores e chegar a um ponto comum. As pessoas que estão preocupadas também fazem uso dessas ferramentas, e isso vai ser importante para o futuro delas. Então, é importante a gente olhar com calma e avaliar principalmente o protótipo desses projetos de lei que estão sendo feitos agora, para aí, sim, a gente ter



SENADO FEDERAL
Secretaria-Geral da Mesa

um modelo que vai inspirar, e vão falar: "Olha só o Brasil, fez um modelo que é referência; conseguiu". É claro que a gente não vai agradar 100% a todos, mas a gente vai chegar a um ponto em comum.

Então, são essas as minhas considerações, Senador.

Obrigado pela oportunidade novamente.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Obrigado, Rodrigo. Aliás, eu peço desculpas, porque eu acabei te colocando numa situação complexa: eu falei que ia ler as perguntas e te passei direto.

Deixe-me ler aqui, agora, com calma, para dar tempo para todo mundo raciocinar também, e também para dar valor àqueles nossos brasileiros que enviaram as perguntas através do e-Cidadania. Aliás, agradeço a audiência pela TV Senado e pelas redes do Senado, e a participação no Portal e-Cidadania, do Senado.

A Deborah, do Distrito Federal, enviou a seguinte pergunta: "Como o projeto equilibra a proteção da privacidade dos indivíduos com o estímulo ao desenvolvimento econômico [...] [por uso de] inteligência artificial?".

O Breno, de Pernambuco: "Como lidar com os modelos de LLM [são os grandes modelos de linguagem, em português] que não têm compromisso com a verdade e frequentemente passam informações erradas?".

O Ney, do Acre: "Quais são os principais fatores a serem considerados na avaliação de riscos em sistemas de inteligência artificial?". De certa forma, eu falei – não tinha lido antes, mas falei – de fatores contribuintes aqui justamente para levantar esse ponto.

A Aline, de Goiás: "Como [...] [o projeto de lei] pretende tratar o uso de [...] [inteligência artificial] para criação de conteúdo pornográfico sem autorização do titular da imagem?". Isso eu vou deixar para os nossos painelistas responderem.

O Marcus, de São Paulo: "Como regular algo de que ainda não temos pleno conhecimento? Até que ponto jurídico podemos usar a tecnologia em questão sem causar danos?".

O Aurélio, de Goiás: "De que forma o PL 2.338/2023 aborda o uso de inteligência artificial na criação artística no que tange à proteção dos direitos autorais?". Isso foi, de certa forma, respondido aqui pelo Rodrigo também.

O Alexandre, de São Paulo: "Como assegurar que a definição de 'risco excessivo' [entre aspas] no PL 2.338/2023 não seja manipulada para beneficiar certos grupos?". Olhe, é criar riscos para produzir mais risco também. E quero lembrar que o ser humano, de forma geral, é o maior gerador de riscos e de problemas que a gente tem. A tecnologia não é exatamente um problema.

A Ana, do Amazonas: "Que tipo de fiscalização pode ser feita para garantir que a [...] [inteligência artificial] não será utilizada de forma [...] [criminoso], como perpetuação de *fake news*?". Lei de *fake news*. Às vezes, eu respondo assim, direto, porque é importante diferenciar as coisas. É a lei de *fake news*, feita para isso.

A Rayane, de São Paulo: "A regularização levará em consideração as leis brasileiras, isto é, privacidade de dados [...] e o uso dos algoritmos?". Sem dúvida. Isso aí está previsto no aspecto como um todo.

A Ana, do Ceará: "O [...] [projeto de lei] oferece algum amparo aos artistas e escritores que têm seus trabalhos usados como material de treinamento para as IAs sem permissão?". É do que a gente está tratando aqui. Na minha opinião, isso também tem que ser tratado de forma, vamos dizer assim, normativa, assim como eleição, assim como aspectos da Anvisa, assim como telecomunicações, assim como outras áreas que precisam ter a sua regulação sob o guarda-chuva da utilização de inteligência



SENADO FEDERAL
Secretaria-Geral da Mesa

artificial, mas é praticamente impossível e muito ineficiente você trazer tudo isso para dentro de uma legislação, porque isso se altera a cada instante.

Comentários do Anderson, de São Paulo: "É insalubre artistas [...] [sofrendo com] suas artes roubadas por inteligência artificial [...] [e com] pessoas [...] [vendendo] imagens geradas por IA".

Do José, de São Paulo: "As ferramentas de IA generativas deveriam ser altamente regulamentadas no Brasil, pois elas geram conteúdo a partir de roubo de terceiros", assim como em outros países. E, aí, a gente tem que sempre lidar com a parte dos outros países. Temos de lembrar que isso aí é transnacional.

Agora, sim, eu vou passar para os outros. Lidas as perguntas e comentários, já com mais tempo para pensar – desculpe aí novamente, Rodrigo (*Risos*) –, eu vou passar agora ao Jefferson de Oliveira Gomes, Diretor de Tecnologia e Inovação da Confederação Nacional da Indústria, para os seus comentários finais, em três minutos, e para responder alguns dos comentários ou perguntas.

Obrigado, Jefferson.

O SR. JEFFERSON DE OLIVEIRA GOMES (Para expor.) – Muito obrigado, Senador.

Eu acho que é mais fácil, para responder ao Breno, à Aline, ao Marcus, ao Aurélio, à Ana, à Rayane, à Ana, ao Anderson, que têm uma preocupação com o compromisso com a verdade, com essa questão de criminalidade, que é melhor sempre dar um exemplo.

Eu tive um aluno de doutorado, que terminou agora o doutorado dele, em que ele aplicou... A gente desenvolveu, na ponta de um robô normal, a ideia para que ele fosse colaborativo, trabalhar com pessoas em volta. A principal causa no Japão, que é um dos principais países que compram robôs, o maior que compra é a China... A principal causa de preocupação na China com os robôs é que as pessoas ficam com medo. O pessoal tem medo de trabalhar com robô colaborativo, já tem medo de trabalhar com o normal, mas tem medo de trabalhar perto de um – imaginem: um robô que veste a pessoa, um robô que fica trabalhando do lado dela. E o que a gente fez? A gente fez um trabalho aqui no Brasil... Ele terminou o doutorado agora e começou há seis anos, há cinco anos ele começou o trabalho. E o trabalho basicamente era verificar – pela dilatação da pupila da pessoa, pela sudorese que ela tinha, pelo batimento cardíaco – o que acontecia com essa pessoa. Conforme ela fosse ficando mais nervosa, o robô começava a trabalhar mais devagarzinho, enquanto que, se a pessoa fosse ficando mais comum – cada pessoa é diferente da outra – ou, então, essa pessoa fosse se transformando, ficando mais confortável, ele começava a trabalhar mais rápido. Para fazer isso, precisou passar pelo comitê de ética da Fapesp. Não tem chance, tem questões. E basicamente o trabalho dele olhou riscos, olhou a aplicação e os riscos. E começou há cinco anos esse trabalho.

Para finalizar, com outro aluno meu, que terminou o doutorado também agora por tempos atrás, a aplicação dele era usando equipamentos de visão para verificar a qualidade de solda de carro, solda de carro. Verifica...

(*Soa a campainha.*)

O SR. JEFFERSON DE OLIVEIRA GOMES – Bota o *machine learning*, verifica como está aquela solda e verifica se o processo anterior foi ruim. Então, faz um *learning* do processo e verifica qual tecnologia. Nós utilizamos oito algoritmos, e, se a gente for olhar hoje quantos algoritmos daqueles que foram aplicados tem, são mais de 35 algoritmos existentes. O processo está ainda em desenvolvimento.

O que a gente tem que cuidar é da lei, compromisso com verdade, proteção ao crime e ponto.



SENADO FEDERAL
Secretaria-Geral da Mesa

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Obrigado, Jefferson de Oliveira, do meu ITA lá também, e Diretor de Tecnologia e Inovação da Confederação Nacional da Indústria.

Eu passo a palavra agora, para as suas considerações finais e para resposta a algumas das perguntas, à nossa cara Ana Bialer, membro do Conselho Nacional de Proteção de Dados e Consultora do Movimento Brasil Competitivo.

Três minutos, Ana.

A SRA. ANA BIALER (Para expor. *Por videoconferência.*) – Senador, muito obrigada.

É difícil em três minutos responder a todas essas perguntas, mas deixe-me tentar trazer talvez uma resposta um pouco mais geral.

Acho que é a Rayane que pergunta sobre a regularização e a consideração das leis brasileiras quando a gente está falando da regulação de IA. Eu acho que isso é muito importante a gente destacar. O marco legal da inteligência artificial não vai existir de uma maneira completamente autônoma e solta no nosso ordenamento jurídico. Toda a inteligência artificial precisa estar adequada aos demais mandamentos e marcos legais que a gente tem no Brasil. Então, muita inteligência artificial trata dados pessoais. Na medida em que a inteligência artificial esteja tratando dados pessoais, ela vai ter que seguir à risca – ou melhor, os agentes de tratamento, os agentes da inteligência artificial terão que seguir à risca – a nossa legislação de proteção a dados pessoais, a LGPD. Até por conta desse contexto, do volume, da relevância dos dados pessoais na IA, é que a gente tem hoje a indicação da ANPD para ser esse grande maestro, essa autoridade central do SIA. Então, é para garantir justamente essa adequação da legislação.

E muitas das perguntas que o senhor trouxe têm a preocupação do direito autoral. Sem dúvida nenhuma, a preocupação do direito autoral na inteligência artificial, principalmente quando a gente fala da IA generativa, é uma preocupação muito relevante, mas a gente tem já em andamento na Casa projeto de lei específico para revisão da nossa Lei de Direitos Autorais, que há tempo sofre críticas por apresentar desafios com relação ao ambiente digital, e aqui antes mesmo da inteligência artificial. Então, parece-me que os desafios seriam melhor discutidos num projeto de lei próprio, mas, ainda assim, os dispositivos específicos de direitos autorais que a gente tem no projeto são meritórios em termos de pontuar a discussão, mas a gente precisa lembrar: não é tudo que é usado num processo de treinamento de inteligência artificial que, de fato, está protegido pelo direito autoral. A discussão é mais complexa, e talvez nós não tenhamos tempo, neste momento, de ter toda a sofisticação que ela, de fato, merece para que nós possamos endereçar isso da melhor maneira possível.

Mas eu volto aqui ao ponto inicial da minha fala. A gente tem que olhar a questão do marco legal de IA como uma parte do ordenamento jurídico nacional e fazer esse exercício, como o Senador trouxe no início: "quais são os crimes que poderiam ser cometidos com o uso de sistemas de IA e que já não estão endereçados de alguma maneira em leis já existentes?". É muito perigoso a gente olhar para a tecnologia, a gente querer regular a tecnologia e a gente fazer esse exercício de uma maneira hermética, desconectando essa tecnologia de todo o arcabouço no qual estamos inseridos. Então, acho que esse ponto é importante pontuar.

Entendo que estamos num momento de reta final no âmbito da CTIA para discussão e votação deste projeto de lei, mas, como falou o Senador Eduardo Gomes na sessão temática outro dia, o caminho ainda é longo. Então, essas ponderações são muito importantes para que a gente possa continuar aprimorando o projeto antes de ele ser aprovado por esta Casa.

Senador, muito obrigada.



SENADO FEDERAL
Secretaria-Geral da Mesa

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Obrigado, Ana. Ana Bialer é membro do Conselho Nacional de Proteção de Dados e Consultora do Movimento Brasil Competitivo.

Eu passo a palavra agora ao Marcelo Almeida, Diretor de Relações Institucionais e Governamentais da Associação Brasileira das Empresas de Software, por três minutos, para suas considerações finais e para responder algumas das perguntas ou em conjunto.

O SR. MARCELO ALMEIDA (Para expor.) – Muito obrigado, Senador.

É importante dizer que a participação social é fundamental, e vejam que a diversidade de temas que está posta aqui revela a complexidade da matéria e, muitas vezes, o desconhecimento daquilo que está sendo debatido, razão que reforça ainda mais o nosso compromisso – e acredito o compromisso desta Casa também – com uma legislação que esteja muito próxima da realidade do cidadão.

Eu fiquei muito curioso com a pergunta do Alexandre, que fala assim: "Como assegurar que a definição de 'risco excessivo' no PL 2.338 [...] não seja manipulada para beneficiar certos grupos?". Essa pergunta me fez refletir. Parece-me que, por trás da pergunta, pode ter uma sugestão legislativa de diferenciação de grupos. Isso é grave. Essa percepção a gente não pode deixar ter. Na minha opinião, isso reforça a tese segundo a qual a questão da alocação de riscos tem que estar tecnicamente dimensionada por uma questão de análise de probabilidade e impacto e não sugerida dentro de um projeto de lei, como está feita pelo PL 2.338.

Tem a questão da Deborah também – eu vou acabar selecionando aqui, Senador, algumas perguntas, sem demérito dos outros questionamentos –, em que ela fala assim: "Como o projeto equilibra a proteção da privacidade dos indivíduos com o estímulo ao desenvolvimento econômico [...] [por meio de uso de inteligência artificial]?". O desenvolvimento econômico é o ponto fundamental que nós precisamos ter como norteador da regra jurídica sobre a qual a gente está se debruçando aqui. Proteger a privacidade dos indivíduos é premissa constitucional, ao mesmo tempo em que o desenvolvimento econômico é. Então, vejam: desde a Constituição, nós já temos orientações segundo as quais nós precisamos obedecer e atingir a regra do meio, do caminho para que a gente possa prestigiar os dois elementos da pergunta da Deborah, que são absolutamente fundamentais.

(Soa a campanha.)

O SR. MARCELO ALMEIDA – O Ney, do Acre, fala assim: "Quais [...] os principais fatores a serem considerados na avaliação de riscos em sistemas de inteligência artificial?". Essa questão acaba sendo circunstancial. Os riscos são tão diferentes quanto os impactos que eles podem vir ou não a ocasionar. Então, mais uma vez, reforço a tese segundo a qual essa matéria não deve estar engessada num PL, num projeto de lei ordinário, razão pela qual eu acredito que isso esteja dentro do que eu chamei de governança legislativa. Isso tem que estar delegado a uma outra questão.

E aí, por último, a Aline fala assim: "Como [...] o PL pretende tratar o uso de [...] [inteligência artificial] para criação de conteúdo pornográfico sem autorização do titular da imagem?". A titulação da imagem já é protegida pela Lei Geral de Proteção de Dados. Então, a gente não pode... Mais uma vez, pegando um gancho também do que a...

(Soa a campanha.)

O SR. MARCELO ALMEIDA – É para terminar, Senador.

Pegando um gancho aqui do que a Bialer falou, existem legislações que cuidam de circunstâncias, como é o caso dos direitos autorais, que está em processo de revisão, e a gente tem que respeitar os



SENADO FEDERAL
Secretaria-Geral da Mesa

seus respectivos lócus legislativos, para que a gente não possa imaginar que a existência da inteligência artificial passa a ser um atrativo regulatório para uma miríade de situações que precisam ser regulamentadas no Brasil. A inteligência artificial não está vocacionada a fazer isso. A gente tem que respeitar esses processos de criação com as devidas legislações postas, para, aí, sim, a gente conseguir fazer uma análise de todas essas circunstâncias, para que a gente possa contribuir para um cenário de segurança jurídica, que, como já foi dito aqui pelo França, está faltando.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Obrigado, Marcelo, pelas colocações muito pertinentes.

Eu queria passar agora, no final aqui, para o Felipe França, Diretor-Executivo do Conselho Digital, para os seus comentários, considerações finais e para responder algumas das perguntas que achar conveniente.

Por favor.

O SR. FELIPE FRANÇA (Para expor.) – Antes de tudo, obrigado, Senador, pela oportunidade.

Acho que a Ana Bialer foi muito feliz quando disse que alguns temas mais específicos precisam ser tratados em leis específicas. E isso foi repetido pelo Marcelo aqui. Então, sobre a parte de direito autoral, talvez a gente não esteja, de fato, maduro para tratar dela. E a gente poderia muito bem ter um projeto de IA sem tratar especificamente desse ponto. Acho que seria mais feliz para a gente ter mais tempo de discutir um tema tão complexo, pois, de fato, é uma escolha política que vai ser tomada. Então, talvez a melhor saída, de fato, seja a gente colocar esse tema para ser discutido em um projeto específico, que trate de direito autoral.

Da mesma maneira, a gente pode falar de diversos outros temas, como desinformação. É um tema sobre o que você tem o PL 2.630, a ser discutido na Câmara.

Nós temos já o marco civil da internet. Por exemplo, sobre a parte de pornografia, se não me engano, o art. 21 do marco civil da internet já fala da possibilidade de remoção de conteúdo a partir da notificação quando é imagem íntima.

Existem já mecanismos, no nosso ordenamento jurídico, para discutir vários pontos.

Assuntos mais específicos, podemos tratar em leis mais específicas.

Aí a gente volta àquela questão: qual é o objetivo da lei? Para a gente, entendendo do objetivo da lei, fica mais fácil definir o escopo.

Nesse ponto, eu queria chegar na pergunta do Marcus, de São Paulo: "Como regular algo de que ainda não temos pleno conhecimento? Até que ponto jurídico podemos usar a tecnologia em questão sem causar danos?".

Essa pergunta é muito boa, porque isso leva a uma grande questão: regular IA, hoje, é como regular a internet em 1980. A gente não sabe o que vai ser a IA daqui a um ano, daqui a cinco anos; dez anos, muito menos. Então, às vezes, é um exercício de humildade não querer resolver tudo de uma vez só.

Nesse ponto, eu acho que o PL 21 da Câmara foi muito razoável ao dizer assim: olha, a gente consegue caminhar até aqui.

Talvez a gente não consiga responder a essas questões, e, para isso, a gente precisa ter um debate mais maduro.

Talvez seja esta a melhor solução: não querer solucionar tudo de uma vez, mas dar um primeiro passo para uma solução normativa...

(Soa a campainha.)



SENADO FEDERAL
Secretaria-Geral da Mesa

O SR. FELIPE FRANÇA – ... garantir a segurança jurídica para o regulador infralegal, definindo um quadrado, até onde ele pode ir, quais são as diretrizes.

Isto é interessante: a gente define esse quadrado, e ele vai ficar mais seguro. O inovador, o empreendedor, quem está criando também vão ficar mais seguros. O usuário vai saber aonde vai. Então, acho que definir esse quadrado, os princípios, as diretrizes, mostra o caminho para onde vai.

Dá para a gente evoluir em cima do texto do 21, adicionar coisas que estavam aqui? Acho que tem muito do caminho do Marcelo: tentar achar esse meio do caminho. E eu acho que o Senador pode ajudar nesse trabalho.

E contem com a gente. É uma oportunidade muito grande debater isto, mas, definitivamente, o texto acho que ainda não está pronto.

E a gente está correndo o risco de querer disputar uma Olimpíada, mas, na verdade, estamos colocando um grilhão no braço de quem está aprendendo a nadar.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) – Excelente. Obrigado, Felipe.

O Felipe França é o Diretor-Executivo do Conselho Digital.

Nós tivemos a exposição de diversos apresentadores, cujo nome eu vou fazer questão de falar e agradecer muito: Patricia Peck, membro do Comitê Nacional de Cibersegurança; Mateus Costa-Ribeiro, fundador da Startup Talisman IA; Rodrigo Scotti, Presidente do Conselho de Administração da Associação Brasileira de Inteligência Artificial; Jefferson de Oliveira Gomes, Diretor de Tecnologia e Inovação, da Confederação Nacional da Indústria; Ana Bialer, membro do Conselho Nacional de Proteção de Dados e consultora do Movimento Brasil Competitivo; Marcelo Almeida, Diretor de Relações Institucionais e Governamentais da Associação Brasileira das Empresas de Software; e Felipe França, Diretor-Executivo do Conselho Digital.

Eu começo agradecendo muito a participação de cada um de vocês pelas informações, pelos comentários, pelas ponderações dentro de cada setor, para mostrar essas visões.

É justamente isso o que a gente precisa aqui.

Eu agradeço também a cada um que nos acompanhou aqui, presencialmente, nesta sala de audiência do Senado Federal, e àqueles que nos acompanharam através da rede Senado, através da TV Senado. Sua participação é essencial.

Isto aqui é uma instituição pública para respeito à democracia, para haver justamente estas discussões.

Aqueles que nos enviaram as perguntas foram: o Marcus, de São Paulo; o Aurélio, de Goiás; o Alexandre, de São Paulo; a Ana, do Amazonas; a Raiane, de São Paulo; a Ana, do Ceará; o Anderson, de São Paulo; o José, de São Paulo; a Débora, do Distrito Federal; o Breno, de Pernambuco; o Ney, do Acre; e a Aline, de Goiás.

Muito obrigado pela participação.

Lembro que não se extingue aí. Nós temos tudo o que foi falado, todas as apresentações ficam no Portal do Senado. Então, vocês podem ir ao Portal consultar e repassarem algumas das falas para entender, porque, às vezes, passam muito rápido aqui. É importante fazer isto: voltarem lá para entender o significado, os dados. Por exemplo, o Jefferson colocou muitos dados ali nas apresentações, são vários números importantes. É bom ir lá e consultarem tudo isso.

Finalmente, eu queria só fazer um comentário nas minhas considerações finais: uma coisa que ficou muito clara, eu acho, durante esta audiência, foi a necessidade da própria audiência. Quando eu



SENADO FEDERAL
Secretaria-Geral da Mesa

pedi a realização dessas audiências, eu fui até criticado, e falaram assim: "Poxa, você quer estender demais a discussão. A gente precisa resolver logo isso aí e tocar em frente".

Eu concordo que é importante ter um processo mais expedito, vamos chamá-lo assim, na área da legislação, no processo legislativo, mas, como eu falei no começo, o fato de nós termos os PLs 21 e 2.338, a Emenda 1, aquela que eu apresentei, e um texto no meio, que é diferente de tudo isso aí, exigia esse tipo de debate, assim como vai exigir das outras duas audiências públicas que nós vamos fazer. E são importantes essa participação e essa discussão para nortear o que o nosso Relator pode fazer e colocar no texto.

Eu sou Vice-Presidente da Comissão, mas obviamente eu não posso fazer o texto; eu não sou o Relator, o Relator é o Senador Eduardo Gomes. Então, é importante que ele tenha essas informações todas, que a gente as coloque para proteção de todo o sistema. O que vai sair daqui, o que vai sair da fábrica certamente não vai ser perfeito, mas é importante que seja o melhor possível. A gente não gosta, por exemplo, que saia o carro com defeito de fábrica e, depois, ter que ficar fazendo *recall* para consertar peças do motor e assim por diante. Então, quanto melhor sair daqui, ideal. Haverá um processo todo pela frente – vai para a Câmara, vai ter votação no Plenário também –, mas é importante. Logicamente, o texto não vai satisfazer a todos, mas não pode, de maneira nenhuma, não satisfazer a ninguém, tem que ter alguma coisa que funcione. Eu coloquei aqui uma série de princípios ou de pilares que eu acho que tem que ser cumpridos e a que uma legislação como essa tem que obedecer. De novo, eu não sou único, não sou eu quem decido tudo, mas é importante.

Eu acho que a gente viu um resultado quase unânime aqui que foi uma certa convergência de fatores. Então, se eu fosse o Relator, eu ia pegar exatamente isso aí e ver: "Pô, está focando em tais pontos aqui. Deixe-me ajustar esses pontos, porque são esses pontos que estão pegando, não estão funcionando bem". E é bom a gente usar o que a inteligência artificial faz, um aprendizado humano sobre isso aí.

O Relator deve vir com um texto novo nesta semana, baseado nessas informações, sei lá, e nas milhares de informações que ele recebeu. Eu vou solicitar a ele que nos apresente o texto um pouco antes, para que a gente tenha um tempo para olhá-lo, porque eu não gosto da situação de você ser apresentado ao texto aqui, como aconteceu na última vez; aí ficar naquela situação: "Não é exatamente isso. De onde saiu isso aí?", e gerar toda essa insegurança. Eu acho importante ter essa apresentação para uma conversa antes.

Alguns pontos aqui que eu vim anotando ao longo do caminho.

Essa legislação precisa ser trabalhada no sentido do desenvolvimento do país, de não restringir o desenvolvimento da pesquisa, o desenvolvimento voltado à inteligência artificial, de não restringir, de jeito nenhum, a utilização da inteligência artificial, de não deixar o medo tomar conta e de usar a lógica para se trabalhar, de se reduzirem as incertezas, se reduzirem as subjetividades, que geram inseguranças de todos os tipos, técnicas e jurídicas. Isso pode ser muito bom para o advogado resolver o problema lá na frente, mas, para o país, isso não é bom, e a gente não quer isso aí de forma alguma.

Pensar que o que nós estamos tendo como maior risco é o excesso de proteção contra riscos, e isso pode ser o maior problema que a gente tenha.

Foram mostrados dados muito interessantes aqui, sobre a quantidade de restrições que nós temos na nossa legislação proposta com a Europa, por exemplo, que já é restritiva demais.

Podemos usar de exemplo os Estados Unidos, para trazer mais negócios. A gente precisa de mais negócios no Brasil; a gente só pode ir para frente, se a gente tiver um setor privado, as empresas tendo



SENADO FEDERAL
Secretaria-Geral da Mesa

sucesso; e não é protecionismo do setor privado, de jeito nenhum. É protecionismo do Brasil, para gerar recursos no Brasil, gerar nota fiscal, gerar empregos.

A gente quer que os nossos brasileiros tenham emprego para trabalhar e não fiquem dependentes de fora o tempo todo.

Simplificar – simplificar! Já existem leis prontas ou sendo preparadas em áreas específicas – a Ana falou muito bem disso. A gente não precisa congregiar tudo dentro de uma lei aqui, o que seria não só ineficiente como também bastante ilógico, porque as coisas vão mudar, e a gente vai ter que ficar mudando isso o tempo todo.

E a gente não deve se esquecer do desenvolvimento das nossas *startups*, não se esquecer da base do ecossistema de que elas precisam.

As empresas grandes têm muito mais facilidade. Você pode colocar uma coisa muito mais restritiva, que elas conseguem se adaptar muito fácil. A *startup* você mata igual uma plantinha pequena, que, para nascer, você tem que cuidar muito bem; aí vai ser o nosso futuro.

Finalmente, a parte de riscos à liberdade.

Eu fico muito preocupado quando eu vejo essa questão de integridade de informação e de você colocar na mão de alguns oráculos a definição do certo e do errado, do bem e do mal, do correto e do incorreto, que isso aqui é verdade, que isso aqui é mentira, que isso aqui é opinião ou que isso aqui é antidemocrático, porque eu sou a democracia, e, se você falar contra mim, está errado. Isso é um absurdo para uma democracia que pretende se manter forte.

Então, eu tenho certeza – vamos chamar assim – de que o Eduardo, nosso Relator, vai prestar atenção nisso, retirar do texto qualquer tipo de coisa que possa trazer risco à nossa liberdade e não colocar na mão de algumas pessoas a definição do que é certo, do que é errado, ou do que é verdade, do que é mentira, porque, se tiver alguém que é dono da verdade, eu ainda estou para conhecer essa pessoa, e, quando a gente partir daqui, provavelmente a gente vai conhecer, mas vai demorar um certo tempo, eu espero, pelo menos para a maioria.

Nada mais havendo a tratar, eu declaro encerrada a presente reunião, novamente com o agradecimento a todos.

Obrigado.

(Iniciada às 14 horas e 42 minutos, a reunião é encerrada às 17 horas e 08 minutos.)