



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

REUNIÃO

19/10/2023 - 8ª - Comissão Temporária Interna sobre Inteligência Artificial no Brasil

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP. Fala da Presidência.) - Boa tarde a todos. É um prazer poder presidir mais uma das nossas audiências públicas a respeito de inteligência artificial, na orientação desse tema aqui no Senado Federal.

Havendo quórum, eu declaro aberta a 8ª Reunião da Comissão Temporária Interna sobre Inteligência Artificial no Brasil, criada pelo Requerimento nº 722, de 2023, com a finalidade de, no prazo de até 120 dias, examinar os projetos concernentes ao relatório final aprovado pela Comissão de Juristas, responsável por subsidiar a elaboração do substitutivo sobre inteligência artificial no Brasil, criada pelo ato do Presidente do Senado Federal nº 4, de 2022, bem como eventuais novos projetos que disciplinem a matéria. *(Pausa.)*

Esta reunião destina-se à realização de audiência pública, com o objetivo de debater a importância da inteligência artificial para a área da saúde, em cumprimento ao Requerimento nº 2, de 2023, da Comissão Temporária de Inteligência Artificial, requerimento esse da minha autoria.

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

Encontra-se presente no Plenário da Comissão o Prof. Edson Amaro Júnior, Professor da Faculdade de Medicina da Universidade de São Paulo, da USP.

Obrigado pela presença.

Em momentos, nós teremos aqui a presença também do Diretor Jurídico da Confederação Nacional da Saúde, Dr. Marcos Vinícius Barros Ottoni.

Encontram-se também, por meio de videoconferência, que vão participar desta nossa audiência: o Dr. Juliano Souza Albuquerque Maranhão, Professor da Faculdade de Direito da Universidade de São Paulo (USP); o Dr. Professor Fernando José Ribeiro Sales, Professor da Universidade Federal de Pernambuco; o Prof. Rodolfo de Carvalho Pacagnella, Professor da Universidade Estadual de Campinas; o Diretor de Tecnologia da empresa Inovia, Tiago José de Carvalho; e o Prof. Wagner Meira Júnior, Professor da Universidade Federal de Minas Gerais (UFMG).

Pessoal, eu gostaria de começar novamente agradecendo a presença não só daqueles que estão aqui, dos nossos apresentadores também por videoconferência e de todas aquelas pessoas que nos acompanham através da TV Senado e das redes.

O tema de inteligência artificial tem várias perspectivas diferentes. Como vocês sabem, nós temos um projeto que foi escrito por uma Comissão de Juristas. Esse projeto está sendo a razão destas análises que nós temos feito aqui nesta Comissão.

Esse projeto precisa ser visto e discutido nas diversas perspectivas existentes não só obviamente na perspectiva jurídica, mas também na perspectiva do desenvolvimento da ciência, da tecnologia associada ao tema, das aplicações disso, das aplicações da inteligência artificial tanto no sentido de desenvolver negócios com a inteligência artificial ou da utilização dentro de negócios. E isso vai ser cada dia mais.

Depois, as implicações dessas aplicações no que concerne às pessoas e como as pessoas vão ser positivamente impactadas por essa tecnologia. E nós sabemos que existem muitas qualidades dessa tecnologia que vão ajudar na qualidade de vida das pessoas.

Hoje, provavelmente, nas nossas discussões, nós vamos ver aqui como ela pode ser utilizada para salvar vidas, para auxiliar o trabalho dos médicos, para auxiliar o trabalho de todos aqueles que trabalham no setor de saúde, o que, em última instância, visa à melhoria das condições de vida e a salvar vidas, mas também nós precisamos analisar os aspectos dos riscos da tecnologia. Toda tecnologia disruptiva tem as suas aplicações positivas e tem os riscos das aplicações negativas. E, quando a gente fala disso, nós temos aí as pessoas como um centro de preocupação em termos dos seus dados, da privacidade dos seus dados, da ética da aplicação, de toda a questão de possibilidade de o aprendizado de máquina ter algum tipo de viés que possa trazer algum tipo de discriminação ou outros problemas. Então, tudo isso tem que ser tratado.

Esta lei é muito importante para o país. Em outros países, leis semelhantes já têm sido aplicadas. Nós estamos observando também esses outros países com relação à efetividade e aos problemas que as leis trouxeram. Isso é importante para que nós tenhamos a nossa lei aqui no "estado da arte", vamos dizer assim: que ela seja flexível o suficiente para permitir o desenvolvimento da tecnologia e dos negócios associados, mas que também ela seja cuidadosa para proteger os direitos das pessoas que vão utilizar, dos brasileiros como um todo.

Passamos à nossa parte prática. Como é que vai acontecer? Nós temos aqui sete apresentadores, ou seja, são sete apresentações, e vamos limitar o tempo de apresentação a dez minutos para cada apresentador. Então, são dez minutos para cada apresentador. Aos apresentadores que estão aqui é fácil de acompanhar isso, porque nós temos o relógio aí na parede, e, então, é fácil. E, faltando um minuto, toca uma campainha, semelhante...

(Soa a campainha.)

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - ... a esta daqui; e, aos 15 segundos finais, aparece um sinal de "faltando 15 segundos". Aqueles que estão remotamente não vão ter esses auxílios aqui com relação ao tempo, e, então, eu peço que vocês prestem atenção aí nos seus relógios, para controlar o tempo; e vocês vão ouvir o aviso de 15 segundos - vai ter uma voz feminina avisando que faltam 15 segundos. Embora seja importante a manutenção do tempo, não cortem, obviamente, o raciocínio pelo tempo ter acabado. Eu vou esperar acabar o raciocínio, eu posso prolongar um pouquinho esse tempo - sem excesso, por favor -, para que a gente tenha a ideia completa sendo apresentada. E, depois das apresentações, nós vamos abrir para debate. Provavelmente, eu mesmo devo trazer algumas questões para alimentar esse debate, assim como as questões que eu espero que a gente receba através do e-Cidadania.

E, para quem estiver aí nos acompanhando, novamente relembro: www.senado.leg.br/ecidadania para mandar as suas perguntas, ou através do número 0800 0612211, para mandar aqui as suas perguntas. É importante essa participação.

Agora, sem mais demoras, eu vou convidar o Prof. Edson Amaro para sentar aqui do meu lado. Daqui a pouco, chega o Prof. Marcos Vinícius também - também o Prof. Marcos Vinícius vai estar aqui no presencial.

E eu já passo a palavra, de imediato, para o Prof. Juliano Souza de Albuquerque Maranhão, que é Professor da Faculdade de Direito da Universidade de São Paulo.

Prof. Juliano, a palavra é sua por dez minutos.

Obrigado.

O SR. JULIANO SOUZA DE ALBUQUERQUE MARANHÃO (Para expor. *Por videoconferência.*) - Obrigado pelo convite, Senador Marcos Pontes.

Na contribuição que eu tive a oportunidade de fazer na terça-feira, eu procurei enfatizar de que modo o advento das inteligências artificiais generativas e das inteligências artificiais em temas fundacionais traz uma transformação no mercado e exige uma revisão do projeto submetido pela Comissão de Juristas, da qual eu participei, no sentido de acréscimo, ou seja, esses sistemas trazem um impacto de transformação no mercado que exigem regras para fomento ao investimento e também para eliminação de gargalos ao investimento em inovação no Brasil. E eu especifiquei como poderiam ser essas regras. Agora, eu gostaria de fazer algumas considerações sobre o que mereceria revisão no sentido de retirar ou atenuar aquilo que está no projeto de lei. Isso diz respeito principalmente ao excessivo detalhamento de regras de governança. O que eu acredito? Na verdade, em todas as iniciativas internacionais de regulação, o caminho tem sido buscar estabelecer regras procedimentais, obrigações de governança para mitigação de riscos por parte dos desenvolvedores e operadores de sistemas de inteligência artificial.

Agora, você tem dois extremos. De um lado, a lei apenas prevê princípios gerais, o que a meu ver é inadequado, porque traz insegurança jurídica, e, de outro lado, tem uma excessiva granularização, um detalhamento de medidas de governança. O advento e a evolução da tecnologia mostram que nós devemos ser cautelosos em detalhar demais, porque esses detalhamentos podem se tornar obsoletos. A minha posição seria no sentido de ter um elenco de medidas de governança que são reconhecidas por organismos internacionais, sobre os quais já há algum consenso, mas sem que haja um detalhamento na própria lei de como essas medidas de governança devem ser implementadas. A própria lei prevê, no art. 30, a possibilidade de entidades setoriais fazerem códigos de autorregulação especificando quais são as melhores práticas dentro de cada setor de atividade. Eu acho que o caminho seria o elenco de medidas de governança e, depois, cada entidade setorial de referência ou reconhecida estabelecer medidas e detalhamentos de governança próprios para o seu setor de aplicação.

Eu gostaria de ilustrar essa minha posição trazendo um exemplo que está presente na lei em relação a regras de governança. Então, deixe-me ver aqui se eu consigo compartilhar. *(Pausa.)*

Eu não estou... Ah, acho que eu não consigo compartilhar, mas eu vou, então, descrever o ponto.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - É possível compartilhar. Estou falando com o nosso pessoal técnico aqui, dá para se compartilhar, sim.

O SR. JULIANO SOUZA DE ALBUQUERQUE MARANHÃO *(Por videoconferência.)* - É, mas acho que é desnecessário, eu prefiro...

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Está aberto lá para ele? *(Pausa.)*

É, acho que está aberto para você o compartilhamento aí agora.

O SR. JULIANO SOUZA DE ALBUQUERQUE MARANHÃO *(Por videoconferência.)* - Sim, mas o meu computador aqui está exigindo uma configuração. Acho melhor eu continuar.

Eu vou dar um exemplo em relação a uma medida de governança, que é o dever de transparência e explicabilidade. É uma medida reconhecida em diversas abordagens de institutos de referência internacionais, como ISO. Agora, a proposta trazida pela Comissão de Juristas - da qual eu participei, ênfase - traz, ali no art. 20, um elenco de medidas de governança para inteligências artificiais de alto risco. Em relação à transparência, está lá no art. 20: "[...] adoção de medidas técnicas para viabilizar a explicabilidade dos resultados dos sistemas de inteligência artificial [...] [ou a explicação de certas decisões], respeitado o sigilo industrial e comercial". É uma previsão genérica, praticamente elenca que deveria haver, nas práticas de governança de operadores ou desenvolvedores de IA, essa preocupação com medidas de transparência e explicabilidade. Ocorre, porém, que existe uma previsão parecida no capítulo sobre direitos, nos arts. 5º, 7º e 8º, que traz um detalhamento do que seriam essas medidas de explicação. Então, o art. 5º fala que as pessoas afetadas têm direito à explicação sobre decisão, recomendação ou previsão. Em seguida, traz-se que as pessoas afetadas teriam direito a receber, previamente à contratação ou utilização, a estarem submetidas ao emprego de IA, uma descrição do sistema, tipos de decisões, o papel também da inteligência artificial e dos humanos, uma explicação do modelo e dos critérios de decisão e das categorias de dados pessoais utilizados. O art. 8º ainda vem adicionalmente e fala na necessidade de explicar a lógica e a racionalidade do sistema, os dados processados e a sua fonte, os critérios para a tomada de decisão e por aí vai, uma série de detalhamentos, o que pode trazer embargos e pode ser inadequado conforme o setor, mecanismos para a pessoa contestar a decisão.

Esses critérios podem ser adequados, por exemplo, para a aplicação de IA em *credit scoring* ou concessão de empréstimos, em que é interessante que aquele que pediu o empréstimo possa ter a explicação sobre os critérios para a avaliação da sua capacidade de adimplência e mecanismos para contestação. E, para alguns modelos, é possível explicar esses critérios.

Agora, se eu for pensar numa aplicação no setor de saúde de um sistema que reconhece manchas em imagens, que analisa imagens de manchas de pele para fazer diagnóstico de câncer, esses sistemas, por exemplo, vão quebrar a imagem em milhões de *pixels*, que não são perceptíveis pelo olho humano, e não há uma definição clara de qual especificamente foi o aspecto da imagem que levaria a uma relação causal com a possibilidade de câncer. Então, são sistemas que podem ter uma acurácia bastante elevada, mas com essa dificuldade de explicação. Nesse caso, sistemas que podem ser extremamente benéficos para o diagnóstico de uma doença grave não poderão ser utilizados, porque não tem uma possibilidade de explicar os critérios efetivos para a tomada de uma decisão ou de uma recomendação diagnóstica?

Um outro exemplo: a especificação dos dados. Cada vez mais, vêm sendo empregados modelos fundacionais. Eu posso ter um modelo fundacional multimodal de descrição de imagens e, para aplicar no campo médico, eu poderia fazer o *fine-tuning* dele para a descrição de radiografias, por exemplo.

Os modelos vão envolver uma quantidade enorme de dados, eu vou aproveitar já um modelo fundacional anterior. Como seria feita essa especificação dos dados relevantes, processada sua fonte? Nós sabemos que, nos modelos fundacionais, no resultado entregue, é difícil você apontar exatamente qual foi a fonte, mas aquilo é construído dentro do conhecimento que é gerado pelo sistema de inteligência artificial.

Esse exemplo eu trago para mostrar como as medidas específicas de explicação de transparência ou, por outro lado, outra medida de governança, de monitoramento do sistema de inteligência artificial podem variar conforme o campo de aplicação: se é no campo financeiro, se é no campo da saúde... E, mesmo no campo da saúde, eu posso ter especificações para medicina diagnóstica que são diversas da indicação de terapias.

O que eu teria a colocar é que aquele detalhamento que vem no capítulo dos direitos em relação às medidas de governança deve ser olhado com cuidado, eventualmente revisto ou até mesmo simplificado ou retirado. Quero dizer que não sou contra uma abordagem de direitos na lei, mas acredito que não se deve confundir uma abordagem de direitos com um elenco de direitos, semelhantes a direitos dos consumidores perante os desenvolvedores e operadores de inteligência artificial, sob pena de criar entraves como esses que eu acabei de trazer para vocês. Então, a minha ponderação seria simplificar a previsão de medidas de governança como obrigações das empresas perante uma autoridade de inteligência artificial e não como direitos com detalhamento excessivo de como devem ser implementadas as medidas de governança.

Lembro também que uma previsão de obrigações no capítulo de direitos migraria a aplicação, o *enforcement* das regras da legislação de inteligência artificial para o Poder Judiciário, o que pode trazer uma série de interpretações conflitantes sobre o que são esses direitos, qual sua extensão, quais as pessoas afetadas, que é outro conceito impreciso que está nos arts. 5º, 7º e 8º, de tal forma que o melhor caminho, repito, deveria ser o elenco das medidas de governança já reconhecidas internacionalmente, acompanhado, então, da autorregulação por entidades setoriais, que podem especificar, por meio de códigos de conduta, qual seria a forma adequada ou quais seriam as melhores práticas setoriais para lidar com a inteligência artificial de um modo responsável, respeitando aquele elenco básico que está na legislação.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado, Dr. Juliano. Parabéns pela apresentação.

Eu já passo agora a palavra diretamente para o Dr. Edson Amaro Júnior, que é Professor Livre-Docente da Faculdade de Medicina da Universidade de São Paulo, aqui conosco.

Professor, o senhor tem dez minutos.

O SR. EDSON AMARO JÚNIOR (Para expor.) - Obrigado, Senador, obrigado a todos pelo convite.

Temos um trabalho longo aqui. Então, vamos ser breves no tempo para sobrar espaço para talvez alguma ponderação.

Primeiramente, é louvável a iniciativa de nós exercermos uma atividade aqui que permite, de maneira pública, levantar algumas questões que, de certa maneira, podem ser melhoradas e aprimoradas no projeto de lei.

Eu vou levantar aqui alguns pontos e ilustrar onde eles se ancoram no projeto de lei. Primeiro, o conceito que já foi colocado pelo Prof. Juliano, muito de acordo com o nosso pensamento, que esse projeto deveria ser pautado por normas principiológicas, de maneira que nós entendamos, sim, que as atitudes e o uso dessas ferramentas possam e devam ser reguladas, porém, não com o nível de especificidade que atualmente está presente na proposta. Eu vou depois entrar em alguns detalhes sobre isso.

Um ponto interessante, por exemplo, no capítulo II, no art. 78, que já foi mencionado pelo Prof. Juliano, só vou chamar atenção das propriedades emergentes que foram observadas em novos sistemas de inteligência artificial, que nós chamamos hoje de modelos fundacionais ou modelos genéricos. Esses modelos têm propriedades que não foram previstas ou planejadas por quem os desenvolveu. Isso acontece, por exemplo, em abril de 2017, os desenvolvedores do ChatGPT, OpenAI, em que eles colocaram, de maneira a declarar sua surpresa, a seguinte colocação: "Uma parte do código era capaz de identificar sentimentos que estavam expressos no texto". O modelo não foi desenvolvido para identificar sentimentos, ele foi desenvolvido para identificar qual a próxima palavra numa frase, isso é bem conhecido, abril de 2017. Nós estamos agora, seis anos depois, discutindo algo que nós percebemos de novo como descoberta. Aqueles que eram técnicos, aqueles que desenvolvem essa ferramenta, perceberam como uma propriedade emergente. Isso hoje é motivo de estudo, isso tem muito das propriedades que a gente ainda não descobriu e, pior, essas propriedades são descobertas apenas com uso de larga escala, não é possível testar isso em laboratório, extensivamente, como se quer, por exemplo, ou como está dito no art. 78 da proposta da lei.

Outra característica importante também bem interessante, claro, isso foi escrito, toda essa proposta, e o texto foi redigido numa época em que nós utilizamos ferramentas de inteligência artificial, que vou chamar aqui de convencionais, em que você treinava um algoritmo com um conjunto de dados, verificava o desempenho desse algoritmo e depois fazia testes

de replicabilidade e também de todas as questões que envolvem, obviamente, a origem, o motor desses algoritmos, que são os dados, os vieses. Isso tudo, de maneira muito pertinente, está colocado no projeto de lei, mas esses sistemas novos não precisam de treinamento extenso, eles são preparados num grande volume de dados, numa arquitetura complexa de difícil conhecimento, portanto, difícil gerar métricas sobre o seu grau de complexidade e risco, mas, quando eles passam a interagir com humanos ou com outros sistemas de inteligência artificial, eles têm uma propriedade que nós chamamos de *Few-Shot Learning*, termo em inglês, que diz basicamente que com uma interação ou duas apenas ele aprende ou demonstra propriedade de aprendizagem.

Como vamos ilustrar isso? Acho que um exemplo bom aconteceu há uma semana, numa convenção em Atlanta, na área da saúde, em que foi colocada a seguinte questão: um paciente com câncer de esôfago, com algumas características, foi perguntado aos principais fornecedores de sistema de IA Generativa ou modelo fundacional, como vocês queiram chamar, LLM, qual seria a classificação daquele tipo de câncer. A gente costuma, em medicina, chamar isso de estadiamento, se o câncer está espalhado ou se está localizado. Então você faz exames, tem algumas características, e o médico tem que saber qual a gravidade. Perguntou-se para a inteligência artificial, modelos generativos, qual seria o estadiamento, ou seja, quão grave estava aquele câncer. Os três produziram respostas diferentes. Isso não é o mais grave, é que, se a mesma pergunta fosse feita depois, qual seria a resposta de cada um deles? Mudou, da mesma maneira.

O que isso traz para a gente? Quando a gente olha de novo a parte de teste de segurança, capítulo III, art. 13, e muito difícil fazer isso com modelos que mudam em minutos. Como você vai constituir um processo que nós precisamos entender melhor, com uma regulação tão estrita? Ela precisa ser mais genérica, ela precisa ser de novo oferecida como uma proposta ou desafio a ser desenvolvido por setores da sociedade com capacidade técnica para tal, para manter o desenvolvimento par a par, *vis-à-vis* com a tecnologia, que avança tão rápido.

Acho também muito importante a gente entender, no art. 32 aqui: será que uma única agência, uma autoridade, pode regular e determinar as especificidades do projeto de lei para nuances de utilização dentro de um mesmo setor, como ilustrado pelo Prof. Juliano, quicá em todos os setores? Então talvez a gente precise transferir para entidades que tenham essa característica técnica, fundamentalmente, nós estamos discutindo ferramentas técnicas cuja compreensão pretende ser um elemento para que nós consigamos ver ou prever seus impactos, eu preciso entender a tecnologia corretamente para prever seus impactos. E é claro que nos assusta muito os exemplos que nós estamos trazendo e têm vindo, a cada instante, de setores diferentes, mostrando como o mau uso pode levar a malefícios.

Porém, eu gostaria de chamar atenção para a área médica, como o bom uso traz benefícios. Nós nunca podemos esquecer que o nosso processo médico hoje é baseado em medicina com lastro em ciência, medicina baseada em evidência. É o que nós aprendemos, o que nós ensinamos. Nós queremos saber se o que a gente pratica no dia a dia é passível de ser reproduzido, entendido que tem um mecanismo biológico, não só por crenças ou ideologias. Como a inteligência artificial ajuda aqui? Porque esse nosso mecanismo, de certa maneira, também não é suficiente para atender uma população tão desigual quanto a do nosso país. Por quê? Vamos pegar um exemplo do paciente com diabetes. Se eu não tenho inteligência artificial, eu vou trabalhar com a ciência que existe hoje, que prediz que a premissa que eu utilize, estudos baseados em projetos de pesquisa com critérios de seleção em que eu seleciono pacientes para entender o mecanismo da doença, mas dificilmente cada paciente que vai no consultório do médico que atende numa UBS vai ter todas as características daquele estudo. O médico tem que se adaptar, a gente tem que fazer essa adaptação.

Hoje a inteligência artificial nos ajuda a fazer isso, mas esses métodos que nós estamos discutindo aqui não podem ser cerceados de maneira a impedir que seu uso traga benefícios, por exemplo, de saber qual a melhor dose para qual paciente, para qual condição e para o lugar em que ele vive. Hoje a gente consegue lidar com a complexidade de informação, de acesso a serviços, que não são só de saúde e impactam tremendamente na nossa capacidade de manter o bem-estar, por exemplo, transporte. Todos esses pontos que nós conhecemos e que impactam a saúde do indivíduo podem ser tratados por inteligência artificial de maneira objetiva, e muitas vezes o médico não está preparado para utilizar isso. Então, quando a gente olha para esse aspecto do ser humano que utiliza essa tecnologia, é importante que esse projeto de lei considere o treinamento das pessoas, a utilização da tecnologia como um dos efeitos, talvez uma das formas de evitar os efeitos maléficos e proporcionar ou potencializar os efeitos benéficos seja a parte educacional do uso de ferramentas de inteligência artificial.

Muitas vezes, quando a gente lê - aqui de novo volto ao art. 32 -, você vê que o aspecto está muito preocupado em características muito específicas, mas às vezes se perde o ponto de vista de quem é que vai fazer, por exemplo, supervisão de IA, de inteligência artificial, é um técnico, é só um médico, como isso será feito? Essas perguntas precisam ser respondidas por um órgão que tenha capacidade de entender a tecnologia, e esse, sim, pode orquestrar, direcionar, fazer interfaces. Por exemplo, uma hipótese, o Ministério da Ciência e Tecnologia poderia, de certa maneira, consultar "Espera aí, esse aqui é o setor da sociedade que envolve saúde" o Ministério da Saúde deve ser consultado para isso.

Essa capacidade de lidar com conhecimentos de aplicação de inteligência artificial e propriedades tecnológicas precisa ser mais bem, eu entendo, trabalhada dentro desse projeto de lei.

Então, como sugestão, terminando o meu tempo aqui, eu gostaria de pensar se não caberia estudar a possibilidade de os órgãos setoriais que possam se incumbir de fiscalizar as adequadas ferramentas, como está previsto, mas não uma agência central. Então, por setor, com uma característica de autorregulação claramente autorizada dentro dos termos, que isso possa ser feito de maneira adequada e que possa atender essas normas. Nós precisamos ter pessoas que conheçam de tecnologia e das especificidades de cada setor, e a capacidade de interface entre esses dois é importante. Promover essa autorregulação com mecanismos que permitam definir...

(Soa a campainha.)

O SR. EDSON AMARO JÚNIOR - ... fiscalizar e manter atualizadas as práticas que garantem os princípios do projeto de lei.

Então, neste último momento, vou frisar um ponto: esse é um processo dinâmico. Nós não sabemos qual é o comportamento do modelo generativo, uma vez que ele é colocado em uso. Citando um artigo científico, recentemente colocado pelos desenvolvedores na Europa, nos Estados Unidos, grandes desenvolvedores dessas tecnologias, que hoje continuam trazendo as novidades, em grande parte, a gente se surpreende com dois pontos.

Primeiro, que se, há pouco tempo, a gente precisava de grandes recursos computacional e financeiro para treinar esses modelos fundacionais, hoje nem tanto. Hoje a gente consegue fazer em computadores de menor capacidade tecnológica, uma vez que nós podemos usar os próprios modelos fundacionais para desenvolver novas soluções. Então o processo é bem dinâmico. Então chamo a atenção para o aspecto dinâmico: monitorar as novas capacidades de inteligência artificial deve ser uma preocupação.

O segundo ponto é que uma vez que nós entendamos que esse processo é complexo e dinâmico, não faz sentido nós tentarmos, mais uma vez...

(Soa a campainha.)

O SR. EDSON AMARO JÚNIOR - ... enrijecer o processo com uma decisão de trazer para um órgão único a capacidade de definir aspectos e detalhar processos com base numa legislação. Talvez precisemos trazer mais pluralidade em quem vai cuidar, se incumbir de fiscalizar o uso adequado das ferramentas de inteligência artificial.

Muito obrigado.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado, Prof. Edson. Parabéns pela apresentação. Também trouxe vários pontos importantes, interessantes aqui para que nós levemos em consideração.

E na sequência, eu passo a palavra ao Prof. Rodolfo de Carvalho. Ele é Professor da Universidade Estadual de Campinas. Prof. Rodolfo, o senhor tem dez minutos para sua apresentação Obrigado.

O SR. RODOLFO DE CARVALHO PACAGNELLA (Para expor. *Por videoconferência.*) - Exmo. Senador Astronauta Marcos Pontes, muito obrigado.

Senador, meu sobrenome é pronunciado de várias formas. Fique tranquilo, porque eu reconheço todas elas.

Fico muito contente pelo convite. É uma oportunidade extremamente importante para a sociedade se manifestar. E eu vou pedir licença para me apoiar em alguns eslaides, para ficar mais fácil que a plateia possa seguir meu raciocínio.

Bom, eu sou professor livre-docente da Faculdade de Ciências Médicas da Unicamp e coordeno a trilha de saúde de uns dos centros de pesquisa aplicada em inteligência artificial, que é o Bios.

Eu queria chamar a atenção para alguns pontos em relação ao projeto de lei, especialmente em algumas palavras que eu destaquei durante essa apresentação. A primeira delas é que esse projeto de lei visa não apenas à implementação responsável, mas ao desenvolvimento de sistemas de inteligência artificial. E esse é um ponto extremamente importante.

O segundo aspecto é que vários pontos são centrais para a área da saúde na apresentação desse projeto de lei, como, por exemplo, a centralidade da pessoa humana, a não discriminação, justiça, equidade e inclusão. São princípios inerentes à área da saúde, que precisam estar então destacados em qualquer ato e em qualquer mecanismo que vá ser implementado na área da saúde.

E eu gostaria então de destacar os pontos relacionamentos aos princípios que estão no projeto de lei. E para isso, eu trago três referências importantes, que eu acho que vale a pena a gente destacar como um arcabouço conceitual para a gente poder fazer uma discussão.

O primeiro deles diz respeito a como as estratégias de saúde e os sistemas de inteligência artificial estão sendo implementados na área da saúde. Existe uma rápida evolução dos sistemas de inteligência artificial, e a disponibilidade de dados agregados e a utilização desses dados nos permitem criar modelos cada vez mais poderosos.

A grande questão é que a maior parte desses modelos que se destacam na literatura não são, de fato, executados na linha de frente da prática clínica, por quê? De certa forma, quando a gente vai trazer um modelo de um sistema de inteligência artificial, a gente não está mudando processo, a gente não está mudando o serviço. E a gente precisa de algo mais, a gente precisa de reestruturação do serviço para que isso seja utilizado.

E mais do que isso, para que isso seja aplicável àquela população, é importante que a população esteja representada nos dados de origem dessa ferramenta. E isso nos leva a uma questão fundamental quando a gente vai pensar em inteligência artificial, a questão dos dados. Não existe inteligência artificial não discriminatória sem dados representativos da população na qual ela vai ser utilizada.

Então qualquer regulamentação que tenha o objetivo de desenvolver essa área precisa desenvolver formas de geração, curadoria de dados que sejam representativos da população brasileira, o que inclui aqui todas as minorias e as populações sub-representadas.

O projeto de lei traz esse termo "não discriminação" em vários artigos. No entanto, há uma necessidade de reforçar como desenvolver e gerar dados e ter uma curadoria adequada desses dados, de forma que isso possa melhorar os sistemas de inteligência artificial.

Uma outra questão que eu vou trazer é uma analogia com relação ao desenvolvimento, à regulação de medicamentos em saúde. Alguns dos princípios de regulação de medicamentos foram estendidos a outras tecnologias médicas e podem incluir aplicações de inteligência artificial em saúde.

Brevemente, muito simplificada, o desenvolvimento de uma medicação envolve um ciclo pré-clínico, com testes *in vitro* em animais, para avaliação de segurança e eficácia; um estudo de fase I, em que eu vou avaliar segurança essencialmente; um estudo de fase II, em que eu vou avaliar a eficácia; um estudo de fase III, em que eu amplio o número de participantes e avalio a eficácia e a segurança. Após isso, eu tenho uma aprovação regulatória para órgãos específicos, no caso de medicamentos, a Anvisa, em que eu tenho que apresentar um dossiê representativo disso, para ter o objetivo de ter um registro dessas vacinas e desses medicamentos. E depois eu tenho uma fase de pós comercialização, que é um monitoramento contínuo de segurança e eficácia.

Eu vou dar um exemplo de um estudo de uma estratégia de um sistema de inteligência artificial que faz esse caminho. Então foi um sistema desenvolvido para avaliação oftalmológica, que ajuda a identificar condições mais graves antes da presença do profissional de saúde. Isso foi testado na vida real, para ver se isso melhorava a qualidade, a produtividade, sem comprometer a qualidade do atendimento. Isso foi mostrado então, por esse estudo, que isso fazia sentido.

Esse seria um estudo que estaria em fase II ou III e que, a partir de então, poderia se solicitar a revisão por órgãos técnicos sobre a possibilidade de utilização na prática clínica de fato. E depois disso, a monitorização contínua de segurança e eficácia desse tipo de modelo. Com isso, o que é que a gente busca nesses sistemas de inteligência artificial aplicados à medicina? A gente tem que buscar eficácia e segurança.

O projeto de lei não especifica de forma muito clara quais aspectos de eficácia e segurança precisam estar descritos. Por outro lado, o projeto de lei traz algumas menções a esta palavra "robustez", que, do ponto de vista técnico, não tem um significado adequado. O que é que eu quero quando eu falo de eficácia? Eu quero que esse sistema apresente velocidade, eficácia, escalabilidade, acurácia, estabilidade, custo e efetividade. E depois eu preciso que haja um processo de contínua avaliação desse sistema. Mas isso não pode estar estancado; isso não pode estar fixo e definido em um projeto de lei que enrijece essas formas como esses temas vão ser avaliados. O Prof. Edson acabou de trazer exemplos muito interessantes desse processo, que é dinâmico e que, portanto, pode ser mais efetivo se houver uma revisão por órgãos técnicos específicos setoriais.

Outra questão importante diz respeito à utilização de estratégias de sistemas de inteligência artificial na saúde, diz respeito à confiabilidade. Os sistemas precisam ser confiáveis, tanto para os usuários finais, quanto para os usuários intermediários, ou seja, para profissionais de saúde, quanto para os pacientes.

E cada vez mais, é mais relevante quando eu trato de questões que dizem respeito a dados sensíveis das pessoas. Para isso, eu preciso de dados, de novo, dados representativos da população, para que eu não tenha um viés demográfico nessa avaliação, nessa modificação contínua desse sistema de inteligência artificial.

Algumas soluções são possíveis, como a validação externa por órgãos reguladores técnicos setoriais, que podem olhar para isso continuamente. Alguns exemplos como, por exemplo, o Japão que introduziu políticas muito interessantes, focando na qualidade, eficácia e segurança dos produtos de inteligência artificial, com o objetivo de equilibrar a automação com a supervisão humana e a decisão humana em algumas questões críticas em saúde.

Então a gente entende que é oportuno que os padrões existentes para os relatórios de modelos de previsão que utilizam sistemas de inteligência artificial estejam atualizados especificamente para incorporar normas aplicáveis para esses fins. E, portanto, eu preciso garantir confiança e validação externa nesses sistemas de inteligência artificial, confiança que vai me dar um monitoramento contínuo de segurança e eficácia. Esse monitoramento vai me permitir avaliar efeitos que só vão ser possíveis de aparecer em escalas maiores, que vão aparecer de 1 para 15 mil, de 1 para 150 mil pessoas, de 1 para 1 milhão de pessoas. O projeto de lei precisa garantir essas formas de monitoramento e o que é decisão humana nesse processo. E, mais do que isso, como deve ser feita a validação externa.

Uma questão também importante e crucial, que já foi colocada aqui também, diz respeito a como os modelos de inteligência artificial respondem, devolvem as respostas. Alguns modelos tipo caixa-preta podem não ser capazes de trazer explicações. Alguns autores indicam a necessidade de a gente desenvolver modelos inerentemente interpretáveis, porque a interpretação desses modelos de caixa-preta é pior do que simplesmente o reconhecimento das respostas em cima desses modelos.

Isso traz um lastro muito importante que diz respeito à responsabilidade de uso do sistema. O projeto de lei traz uma responsabilidade restrita aos desenvolvedores, mas existem responsabilidades solidárias que envolvem os administradores hospitalares, a equipe de saúde, os consultores em tecnologia, que precisam ser consideradas nesse momento.

Então, a questão da transparência é fundamental na utilização dos sistemas de inteligência artificiais em saúde. A responsabilização também deve ser considerada, mas enfrentando esses desafios que a gente tem com a ampliação dos responsáveis por esses ciclos de inteligência artificial, inclusive os sistemas que utilizam código aberto.

Por fim, é muito importante que a gente estabeleça uma estratégia nacional que é muito maior do que um plano regulatório. Uma estratégia nacional envolve que a gente quer buscar desenvolvimento e que se tem que buscar estímulos de geração de dados representativos da nossa população, estratégias que busquem eficácia, segurança e confiança nos modelos de inteligência artificial.

Agradeço a possibilidade dessa minha fala. E esse é o recado que gostaria de dar.

Muito obrigado.

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

Passo agora a palavra ao Prof. Wagner Meira Júnior, Professor Titular do Departamento de Ciência da Computação da Universidade Federal de Minas Gerais.

Prof. Wagner, o senhor tem 10 minutos.

O SR. WAGNER MEIRA JÚNIOR (Para expor. *Por videoconferência.*) - Boa tarde a todos. Aproveito para agradecer o convite. Estou muito honrado por estar aqui tentando dar uma contribuição para essa discussão, que é extremamente relevante.

Então, como o Senador mencionou, eu sou Professor de Ciência da Computação. Então, a minha abordagem vai ser eminentemente técnica nesse processo.

Agradeço muito a quem me antecedeu, porque já antecipou alguns dos tópicos e temas e embasou, inclusive, alguns dos argumentos que eu quero trazer aqui.

Um aspecto que foi colocado agora por último é que, apesar de a saúde ter sido sempre um cenário de aplicação de inteligência artificial, como por exemplo o projeto, da década de 70, de Stanford, isso nunca chegou, realmente, a virar uma prática clínica. Uma das perguntas é exatamente por que isso aconteceu, mesmo após décadas de pesquisa no assunto? Talvez porque o processo diagnóstico é muito mais subjetivo do que um conjunto de regras, que era o que se propunha na época. E o Edson e o Rodolfo foram muito hábeis em discutir isso.

Quando nós colocamos por que a IA em saúde é desafiadora, é importante realçar que, por trás de toda inteligência artificial nós temos um modelo que é uma abstração da realidade e quaisquer problemas de dados, como foi muito bem colocado por eles, limita a capacidade e o alcance desses modelos.

Nós temos vários exemplos prototipais de modelos desenvolvidos que foram capazes de replicar a capacidade de um médico na sua avaliação, por exemplo, de exames, que são os mais populares. Pessoalmente, eu tive uma experiência em que a gente foi capaz de alcançar a capacidade de um cardiologista na avaliação de um ECG. Mas, ao mesmo tempo em

que ele acertava para um conjunto significativo de doenças, ainda tinha em torno de 12%, 13% que não conseguia acertar. E é aí, exatamente, que vem o problema, um desafio clínico muito grande.

Uma outra coisa que eu acho que é um caso interessante - e que, talvez, seja tão natural para o Edson e para o Rodolfo que eles não trouxeram e que, para mim, muito assustou - é que o normal nem sempre é bem definido em saúde. Então, o normal é algo que tem uma ausência de doença, de uma patologia, de um diagnóstico. E isso, do ponto de vista de um modelo, é algo assim assustador, porque ele nunca vai saber direito. Esses modelos de alta dimensionalidade, como é o caso das tecnologias generativas e das redes de aprendizado profundo, facilmente podem - que seria o termo técnico - sair pela tangente ao se depararem com uma situação normal, mas desconhecida. Muitas vezes, mesmo o volume de dados não é suficiente para resolver esse tipo de problema.

Eu conversava, nesses dias, com um professor do Centro Nacional de Vacinas aqui da UFMG, e ele me relatou que o vírus do HIV muda mil vezes mais do que o vírus do covid. E isso, até hoje, tem impedido a construção de uma vacina que seja efetiva para a Aids. Então, há questões intrínsecas à saúde que podem ser basicamente impraticáveis para qualquer sistema computacional.

Por outro lado, esses modelos, considerando um conceito do que eles capturam de forma latente, podem ser muito poderosos. Então, nós temos um exemplo aqui, naquele mesmo modelo de diagnóstico de eletrocardiograma que eu mencionei, que também foi capaz de estimar com uma precisão - acima de qualquer cardiologista - a idade cardíaca do paciente, inclusive interpretando o quanto que a pessoa teria de sobrevida e a sua expectativa de vida.

Esse aspecto é interessante porque traz aqui uma analogia, dentro do meu pouco conhecimento, daqueles fármacos que servem para outras doenças que não a original. O mesmo acontece na inteligência artificial. E aí - como nós vimos, já há alguma discussão - tentar delimitar claramente para que uma dada tecnologia vai servir ou até onde ela pode ir é algo muito complicado do ponto de vista técnico.

Então, eu acho que seria mais interessante se a gente buscasse formas de mostrar e de permitir que os profissionais, em particular no nosso caso aqui os profissionais de saúde, utilizassem algo que é inquestionável nesses modelos de inteligência artificial, que é a possibilidade de mostrar e comparar cenários alternativos, de abrir possibilidades e de verificar outros diagnósticos ou outras situações que originalmente poderiam não ter sido levantadas pelo profissional de saúde.

Nós temos um problema bastante sério, que o Rodolfo e o Edson colocaram com muita propriedade, que várias das propriedades, dos princípios que são colocados, por exemplo, no projeto de lei, são... Eu concordo que eles são chaves como, por exemplo, transparência, auditabilidade, confiabilidade, mas, ao mesmo tempo em que são chaves, são dinâmicos. Tem uma questão que é bastante interessante. É que, quando nós falamos, por exemplo, dessas duas que eu mencionei - transparência e auditabilidade - elas não são características intrínsecas do sistema, são características que dependem do usuário do sistema, no caso, do profissional de saúde. Então, a transparência vai ser tão melhor quanto mais o profissional de saúde confiar naquele sistema; a auditabilidade vai ser tão mais efetiva quanto mais formos capazes - o que nem sempre é possível, considerando modelos opacos - de reproduzir os conceitos e a base nos quais esses sistemas foram baseados.

Então, existe uma questão - que eu vejo como uma separação e que apenas alguns pontos do projeto de lei menciona - que tira um conceito que tem sido falado, o do *human in the loop*, então, talvez, pelo menos para as aplicações de saúde, em particular para as aplicações de alto risco, tira-se o humano do *loop*, o que eu não sei nem se é possível. Eu vou voltar nisso daqui a pouco.

Nesse caso, a gente pode ainda realçar, corroborando o que o Juliano também colocou, que vários dos princípios e dos detalhes que são apresentados no projeto de lei ainda são objeto de pesquisa. A própria transparência... Nós ainda temos um volume de artigos significativos sobre transparência, inclusive sobre esse fenômeno que foi reportado por ele ou pelo Rodolfo, salvo engano, de que às vezes os modelos transparentes têm um desempenho aquém dos modelos opacos a todas as perguntas de confiabilidade e robustez.

Uma coisa que eu achei muito peculiar e que mostra como essa questão ainda é rica é que, há alguns dias, um psiquiatra me abordou e falou que o chatGPT não alucina, ele tem delírios, porque para alucinar ele teria que ter alguma outra característica - o que sinceramente eu não me lembro e nem tenho capacidade de falar. Mas isso mostra como nós ainda estamos num cenário em que, ao mesmo tempo, é uma grande oportunidade, mas é uma questão muito aberta. Como bem colocado, a emergência, por exemplo, das arquiteturas generativas, nos mostrou como rapidamente uma legislação que seja muito detalhada e que se baseie na tecnologia atual pode se tornar obsoleta.

Um aspecto que eu quero comentar aqui, para terminar quase a minha fala, é que quando nós olhamos, em particular, os sistemas de saúde, eu entendi que praticamente todos são sistemas de alto risco. E isso já é extremamente regulado.

Eu reforço aqui o meu ponto porque eu não consigo ver esses sistemas de alto risco sem que haja um operador que já é regulado, acompanhado, extremamente escrutinado na sua atuação e que vá utilizar... Que esse sistema de alto risco vá ser utilizado de forma autônoma ou de forma genérica, etc.

Então, essa é uma separação que talvez seja importante.

Uma outra coisa que me chamou a atenção, finalizando aqui em relação ao projeto de lei, é que várias das dimensões ainda são apenas qualificadas em termos de adequadas, razoáveis, apropriadas, e isso, do ponto de vista de quantificação e de, por exemplo, avaliação de risco, que é uma palavra utilizada, me parece bastante complicado. Eu, como técnico, não tenho clareza de como quantificar várias daquelas dimensões e, mais ainda, de se a minha quantificação seria relevante, significativa e efetiva.

Eu acho que o Rodolfo foi muito feliz quando falou de eficácia e segurança, e talvez isso exija que essas aplicações de alto risco tenham que ser tratadas sob a perspectiva das respectivas áreas, porque, sendo de alto risco, elas têm que, de alguma forma, realizar esse mapeamento com a prática e os estamentos que são próprios daquela área.

Por fim, nós temos, por exemplo, quando nós falamos de avaliação de impacto algorítmico, eu achei a iniciativa fantástica, mas, assim, é algo que é estado da arte da pesquisa. Então, talvez, se nós tivéssemos a capacidade de distinguir aquilo que ainda é, por mais relevante, algo a ser discutido, a ser estabelecido e algo que já tem uma maior construção, uma coisa mais bem colocada pelas várias comunidades envolvidas, eu acho que seria mais adequado.

Chamou-me a atenção, por exemplo, que lá na Seção III existe um tópico sobre gravidade das consequências adversas. Então, notem que, neste cenário em que a gente mal consegue prever para que as tecnologias vão ser utilizadas, nós sermos capazes de buscar todas as consequências adversas eu acho que é algo bem longe da nossa realidade.

Agradeço novamente a oportunidade e devolvo a palavra ao Senador.

Muito obrigado.

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

Aliás, um abraço aí para o pessoal, também, do Centro de Vacinas. Parabéns pelo trabalho que fazem também!

Gostaria, então, de passar a palavra para o Prof. Fernando José Ribeiro Sales, que é Professor do Departamento de Engenharia Biomédica da Universidade Federal de Pernambuco.

Prof. Fernando, tem dez minutos.

O SR. FERNANDO JOSÉ RIBEIRO SALES (Para expor. *Por videoconferência.*) - Obrigado, Senador Marcos.

É um prazer estar aqui com vocês contribuindo com essa discussão, com esse processo rico.

Parabéns ao Senado Federal por permitir esse processo de consulta ampla numa questão tão fundamental que é a questão de discutir a regulamentação da IA no Brasil, a inteligência artificial, pelos seus aspectos fundamentais em diferentes áreas de conhecimento!

Como o nosso colega já falou aqui, é muito rico estar aqui e já pegar... Vários pontos que eu queria colocar já foram abordados aqui.

Ultimamente, eu tenho me aproximado mais da área de inovação, e, toda vez em que a gente vai discutir questões ligadas à inovação, a gente tenta entender qual é o problema que nós queremos resolver. De certa forma, a gente tem visto aí a inteligência artificial crescendo, a gente vivenciou, no nosso dia a dia, todos esses pontos já colocados e o projeto de lei aparece como uma solução. Mas a solução para resolver qual problema? O que regulamentar? Do que a gente precisa?

Então, tem aí uma série de pontos que já foram colocados e bem colocados pelo Juliano, pelo Rodolfo, pelo Edson e por todos aqueles. E, se eu pudesse sintetizar de uma outra forma, será que nós temos maturidade para regulamentar tudo o que está, entre aspas, "previsto" na proposta de projeto de lei? O porquê é muito claro. Agora, sobre o como e o quê, eu acho que ainda nós não temos clareza. Será que um projeto de lei ordinária é um instrumento adequado para regulamentar tudo?

Eu pergunto isso com a visão, talvez, um pouco mais, Senador Marcos, de engenheiro, como o senhor. A gente gosta muito de perguntar sobre o porquê. Será que eu estou usando a solução adequada para o problema? Muitas vezes, nós gastamos muita energia focando a solução e não o problema. Então, a gente está se apegando a isso aqui.

De fato, ter uma lei ordinária é alguma coisa importante, principalmente naquela parte, na minha visão, inicial, de garantir, talvez, direitos universais em relação aos impactos que a IA e as ferramentas que usam a inteligência artificial podem criar. Mas, analisando isso numa visão não jurista - o Juliano colocou, tem uma Comissão de Juristas -, tem uma série de questões com que eu me preocupo em relação à viabilidade, de a gente conseguir, de fato, implementar tudo o que está ali previsto.

Como o Edson e os outros colegas colocaram, é um processo bastante dinâmico. Nós sabemos que a tecnologia gira num ciclo muito mais rápido que as pessoas, que giram num ciclo mais rápido que as entidades, que giram mais rapidamente que as políticas públicas - quanto mais o processo de regulamentação!

Então, será que regulamentar, definir em uma lei ordinária questões, entre aspas, "num nível mais baixo de governança", detalhes de operação, é a melhor forma? É o que nós queremos? O Juliano colocou isso muito bem.

A gente tem preocupações, é natural. São naturais essas preocupações, principalmente porque a gente pensa... Hoje em dia, a gente usa o ChatGPT e outras ferramentas como essas e vê muita coisa sendo gerada, como vários exemplos foram colocados pelos colegas aqui, mas a questão toda é que a gente está num nível hoje em que a gente tem sistemas, pensando para a saúde, sistemas de auxílio à decisão em que os profissionais ainda são responsáveis por emitirem, tomarem essa decisão.

Então, é claro que a gente precisa investir na formação, criar condições de ter melhores profissionais, mas a decisão ainda é do profissional médico, é do profissional de enfermagem, do farmacêutico, de todas as áreas. Mas ter ferramentas baseadas em inteligência artificial pode aumentar bastante a escala, ampliar o impacto que nós queremos.

Então, eu acho que, para a gente avaliar essa questão da maturidade, é fundamental que a gente conheça ainda mais as questões. E aí, quando a gente vai para a proposta escrita... Acho que foi o Rodolfo, ele começou com o art. 1º. A primeira palavra que ele destacou foi "desenvolvimento". Do jeito que está escrito ali, eu fico com receio de entender que, do jeito que está ali, o projeto de lei pode atrapalhar, inclusive, o projeto de vários pesquisadores que estão com pesquisas em andamento, porque uma coisa é o desenvolvimento da solução ou regulamentação para aprovação em uso lá na frente. Mas, se eu tenho que prever todas aquelas questões até para prever o desenvolvimento de potenciais soluções, talvez eu esteja travando um processo importante de experimentação.

Então, a proposta de projeto de lei, se pensarmos na hierarquia das normas, já tem uma série de outros órgãos, por exemplo, entidades de classe, que têm seus códigos de ética, códigos de regulação, que, de certa forma, já endereçam, por exemplo, questões em relação à segurança de informação, questões em relação ao uso de dados em saúde - por exemplo, o Código de Ética Médica; você tem regulamentação que diz o que é um prontuário eletrônico, o que você pode fazer com o dado que está ali. Como é que isso se relaciona com essa regulamentação que vai ser proposta? O Edson colocou muito bem: jogar tudo isso dentro de uma regulação única, uma regulamentação única de uma autoridade competente única, que, na minha visão, não está bem descrito o que será essa autoridade? Ela vai tratar de alguns, ter autonomia para atacar algumas questões, mas ela vai ser uma agência reguladora tipo a Anvisa, tipo a ANS, tipo a Anac? Mas vai ter matérias, quando nós pensarmos na aplicação da IA em cenários específicos? Como é que a gente vai lidar com o conflito de regulamentações de um mesmo nível, ou seja, de uma Anvisa com uma agência regulatória de IA - não é isso?

Então, tem uma série de questões também ligadas ao exercício profissional. A gente não pode esquecer, por exemplo, a Anvisa... E o Rodolfo trouxe, fez uma aula muito boa na questão de como é que tem o processo, por exemplo, de avaliação quando você vai desenvolver um fármaco. Você tem processos hoje que permitem, vamos dizer, têm clareza para dizer: "Eu quero um equipamento médico. Eu quero aprová-lo para uso". Existem processos para isso também. E tem processo para *software* como dispositivo médico, que vêm numa tentativa de acompanhar as discussões no mundo, envolvendo esse *software* como dispositivo médico, usando IA também.

Então, é fundamental nós viabilizarmos meios de garantir a experimentação. O Juliano falou muito bem disso hoje. Falou no dia 17, salvo engano. Ele destacou isso bastante, a questão de que nós temos de pensar também num fomento para potencializar o desenvolvimento da inteligência artificial no Brasil, principalmente como uma estratégia nacional. Os outros colegas colocaram isso aqui.

Então, estratégia como processo de ter aspirações para chegar a resultados. E aí eu vou pegar um pouquinho na seção que eu mais gosto ali, que é a parte de fomentar a inovação, na Seção III. Quando nós falamos de inovação, tem várias definições, mas a que eu mais gosto é aquela atribuída ao Peter Drucker, que fala aqui: Inovação é quando você tem mudança de comportamento dos atores ou dos agentes em um mercado por meio de consumo de produtos, serviços e outras questões.

Nós queremos mudar comportamento. Nós queremos construir produtos, serviços, instrumentos que vão melhorar a saúde, mudar o comportamento. A gente tem desafios enormes para lidar num país como o Brasil. Então, é fundamental que nós tenhamos ambientes de experimentação.

Uma coisa muito boa desse projeto de lei é a previsão da criação de *sandboxes* regulatórios. Talvez, o como criar, quais vão ser as diretrizes e se isso deve estar no projeto de lei ordinária, eu não sei, mas eu acho que é fundamental a previsão desses ambientes, que vêm sendo discutidos em vários outros lugares. Mas, para tudo isso, tem que ter uma grande estratégia, uma visão de futuro. A IA pode ser, num país como o nosso, uma grande ferramenta para habilitar e potencializar muitas das riquezas que nós temos.

Então, se a gente for fazer um paralelo, por exemplo, hoje nós fazemos parte do seleto grupo no mundo que consegue produzir aviões, porque, lá atrás, décadas atrás, teve-se uma visão de futuro e aí a gente viu, criou-se uma estratégia, um processo, primeiro, uma escola para formar engenheiros aeronáuticos e, hoje, temos uma empresa que faz aviões e faz lá toda a parte de mobilidade e que é referência no mundo. Por que a gente não pode fazer isso no Brasil?

A gente está falando muito sobre risco. Isso é importante, mas também temos que olhar os benefícios, como o Edson colocou.

Se a gente pensar nos impactos de redução de tempo, de custo, quando a gente pensar em escala... Se a gente for pensar na questão, por exemplo, de novos fármacos, é fato, a gente reduz bastante o horizonte de busca, a gente consegue otimizar bastante esse período.

Se a gente tiver que ter toda aquela previsão dos efeitos, que estão previstos, não me lembro agora do artigo, isso pode ser muito crítico.

Eu vou finalizar aqui.

Toda a questão de bioinformática, a questão dos vieses que foram colocados. Uma das grandes tendências para reduzir viés é usar geração de imagens, para reduzir desbalanceamento.

Acho que a gente está focando muito em regulamentação, quando é o momento de focarmos em responsabilização, criar os meios para isso ficar claro. Lembrando que, com inteligência artificial, a gente pode dar escala, aumentar a resolutividade.

Nós temos um SUS que, em princípio, é maravilhoso, num país com um tamanho... Mas nenhum país do mundo, com mais de cem milhões de habitantes e com as dimensões que nós temos, tem um sistema de saúde baseado em princípios tão bonitos, em essência, mas que são muito difíceis de realizar. E a gente pode estar amarrando, limitando o potencial de entregar essa escala, de uma forma até sem querer, muito sem maturidade suficiente para fazer isso.

Muito obrigado.

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

Passo agora a palavra, por dez minutos, ao Sr. Tiago José de Carvalho, Diretor de Tecnologia da empresa Inovia.

O SR. TIAGO JOSÉ DE CARVALHO (Para expor. *Por videoconferência.*) - Em primeiro lugar, obrigado pelo convite. Obrigado, Senador. Obrigado ao Senado Federal pela oportunidade de estar aqui falando um pouquinho da nossa visão sobre esse projeto de lei que visa regulamentar o uso de IA, não só na saúde, mas de forma geral no Brasil.

Isso é uma das primeiras coisas que eu gostaria de pontuar, de trazer; que, apesar de nós estarmos discutindo aqui nesta Comissão, nesta reunião, especificamente, o uso de IA em saúde, nós estamos discutindo um projeto de lei que é focado no uso de IA como um todo.

Então eu estou focando no uso de IA na saúde, mas também estou focando no uso de IA em sistemas distribuição de energia, em sistemas de distribuição de saneamento, e etc., em outros sistemas críticos, não só para a evolução científica, mas também para a evolução do Brasil como uma sociedade.

Tem vários pontos, como o Fernando acabou de mencionar, que IA pode viabilizar, pode potencializar muito.

Só que, lendo o projeto de lei, como está escrito hoje, eu tenho uma espécie de receio se esse projeto de lei não está tentando delimitar, definir as coisas com uma visão estática do que acontece hoje, sem olhar para um futuro.

Por exemplo, o art. 17 do projeto de lei define, discorre, sobre sistemas de alto risco. E, dentro desse artigo, todos os sistemas que envolvam saúde, energia e saneamento são de alto risco. Mas aí, como cidadão e como gerador de soluções de inteligência artificial eu me pergunto: realmente, todo o sistema de energia, de saúde ou de saneamento é realmente um sistema de alto risco?

Você pode, por exemplo, limitar um sistema de IA que faz a leitura de um hidrômetro e automatizar isso para dizer que é um sistema de alto risco?

Então, a minha crítica aqui vai, não ao que está sendo dito, mas a como está sendo dito.

Você classificar tudo como sendo a mesma coisa, como sendo tudo sistema de inteligência artificial, é um risco muito grande, porque cada sistema vai ter sua especificidade, cada sistema vai ter sua particularidade e, como os colegas pontuaram muito bem, cada sistema deveria ser regulado por órgãos que são específicos, que têm um conhecimento para aquilo; que vão conseguir colaborar de uma forma mais precisa naquele assunto.

Outro artigo, também, que me chamou a atenção, e que mostra para mim que o projeto de lei busca ser generalista, busca cobrir o maior número de casos possíveis, mas que pode trazer problemas para a sua implantação, trazer problemas para a sua aplicação, de fato, é o art. 20.

No art. 20 a gente tem um ponto onde está escrito assim: avaliação dos dados com medidas apropriadas de controle de vieses cognitivos humanos que possam afetar a coleta e a organização dos dados - novamente. Mas, assim, quem vai fazer essa coleta? Como é que você vai viabilizar esse controle? Você vai deixar isso na mão de uma única agência que, muitas vezes - e a gente sabe que isso acontece no Brasil, pelo excesso de coisas que se tem para fazer - vai formar um gargalo ali e você vai travar a inovação, você vai travar que esses sistemas cheguem ao público de uma forma responsável, o que é muito importante, mas também de uma forma rápida?

Então, assim, o projeto de lei tem algumas palavras ali dentro, alguns termos chaves que são repetidos. Os colegas já destacaram isso aqui, ao longo de todas as falas, que eles são repetidos sistematicamente, mas que, dependendo da forma como se interpreta, aquilo pode, não só responsabilizar, não só incentivar a responsabilização das pessoas no desenvolvimento de sistemas de IA, que é o que se está buscando aqui, mas, ao invés disso, pode travar o desenvolvimento desses tipos de tecnologia, desse tipo solução.

Não vou me alongar muito na minha fala, porque os colegas cobriram muita coisa, mas o último ponto que eu gostaria de falar aqui, de trazer para a discussão, é realmente sobre os ambientes de *sandbox*.

Já fui professor por um bom tempo, hoje sou Diretor da Inovia aqui, que é uma empresa focada em trazer inovações para o mercado e, quando a gente fala de ambiente *sandbox*, eu acho que a definição, dentro do próprio projeto de lei, está muito... O projeto focou em outras partes, em ser muito detalhista, e nessa parte que está falando da inovação, de como a gente pode viabilizar, de como a gente pode fomentar o uso da tecnologia de inteligência artificial para melhorar a vida do cidadão, ele peca por, simplesmente, não dar a atenção merecida, na minha opinião.

Então, eu acho que é um tópico que deveria ser melhor discutido, melhor elaborado e, como o Fernando colocou magistralmente aqui, é um tópico que pode fazer a diferença para a vida das pessoas. Pode potencializar um SUS, pode melhorar o atendimento, só que ele tem que ser melhor elaborado e discutido.

Muito mais do que colaborar, eu acho que os colegas já falaram de forma magnífica, eu só me preocupo com a atenção e com a forma como o projeto está escrito.

Então, acho que todas as manifestações dos colegas aqui fazem muito sentido e deveriam ser levadas em consideração.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado Tiago. Parabéns também pelas colocações.

Neste momento, eu gostaria de aproveitar para ler algumas das colaborações que nós tivemos através do e-Cidadania.

Aliás, lembro ao pessoal que está nos assistindo, pela televisão do Senado ou pelas redes sociais, que ainda há tempo de participar, através do *website* do Senado www.senado.leg.br/ecidadania, para enviar suas perguntas ou colocações sobre o tema, assim como pelo telefone 0800 0612211.

Vou ler aqui algumas dessas perguntas que já chegaram.

Daniela Pessoa, de Sergipe: "A [...] [inteligência artificial] já está sendo utilizada na área de saúde do Brasil?".

Eu estou lendo aqui para os nossos apresentadores também já terem essa... De repente, se alguém quiser responder, depois, fique à vontade.

Primeiro, eu faço uma passagem, para que possamos responder.

Andrei Luan, do Rio Grande do Sul: "O [...] [Sistema Único de Saúde] está preparado para [...] a [...] [inteligência artificial]?".

Rafael Salvatore, do Amazonas: "Como fica a questão do direito ao livre esclarecimento do paciente em relação à inteligência artificial?".

Acho que foi falado aqui também, bastante, sobre a transparência e outros pontos como esse.

Daniel Diniz, da Bahia: "Como será a regulamentação e a fiscalização das atividades assistidas pelas ferramentas de [...] [inteligência artificial]?".

Então são essas quatro perguntas que eu gostaria de colocar para os nossos debatedores, os nossos apresentadores, para que possam responder. Então fica aberto. Quem quiser responder, é só fazer um sinal aí, levantar a mão, a qualquer uma das perguntas. Fiquem à vontade para entrar no tema.

Está aberta a palavra aos nossos debatedores.

Está aqui o Rodolfo.

Rodolfo, tem a palavra.

O SR. RODOLFO DE CARVALHO PACAGNELLA (Para expor. *Por videoconferência.*) - Senador, muito obrigado.

Eu posso começar essa conversa, porque eu acho que dá para dinamizar e compartilhar a palavra...

Eu começo pela provocação da cidadã Daniela Pessoa, se a IA está sendo utilizada na área de saúde no Brasil.

A resposta é sim. A inteligência artificial está sendo utilizada em situações que a gente nem percebe.

Esse é um uso muito interessante, porque ela está em sistemas de apoio à decisão clínica, ela está em sistemas de apoio a métodos de avaliação radiológica, de diagnóstico radiológico, ela está em sistemas de gerenciamento hospitalar, ela está na comunicação com os pacientes, em *chatbots* de comunicação...

Então, a inteligência artificial invade várias formas, está, sim, já permeando grande parte do atendimento privado e, inclusive, do atendimento público - o SUS tem uma série de ferramentas que se utiliza de inteligência artificial - e vai permear, cada vez mais, o atendimento, porque melhora e amplia as possibilidades, como foi colocado, amplia acesso.

A utilização de ferramentas de sistemas de inteligência artificial tem mostrado possibilidades, potências muito grandes, especialmente para o SUS.

Eu passo a palavra aos colegas.

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

Vejo que o Prof. Fernando quer fazer um comentário também. Por favor, professor.

O SR. FERNANDO JOSÉ RIBEIRO SALES (Para expor. *Por videoconferência.*) - Eu queria continuar, como o Rodolfo falou, e o Edson, que está aí ao lado, pode até complementar o exemplo de que eu vou falar, porque está relacionado a ele.

A gente usa, e tem muita coisa já pronta para usar, mas que ainda está em fase, vamos dizer assim, "não usa porque ainda estão esperando". Por exemplo: um dos projetos que está aí feito, capitaneado pelo Hospital Albert Einstein, com cujo gestor eu tive contato, recentemente, que é o Dr. Pedro Vieira.

O programa é chamado - eu esqueci o nome... - Maria. Eles construíram um sistema... Vejam, há quase seis anos, eles começaram um projeto e construíram um sistema que consegue ler qualquer imagem médica e criar uma grande base de dados, priorizada em três grandes categorias - corrija-me, depois, Edson, se eu estiver errado -: imagens acho que de tomografia, para questões relacionadas ao diagnóstico de AVC e acho que alguma coisa relacionada a trauma; lesões de pele, para diagnóstico de melanoma; e também raio-X de tórax, para questões ligadas a tuberculose e outras.

São modalidades de imagem que estão relacionadas a dores reais. A gente tem uma série de pessoas, hoje, que sofrem, porque demoram muito tempo e não conseguem ter resolutividade. Então, eles construíram um piloto para validar, usar imagens e ter essa plataforma para rodar algoritmos de *machine learning* que têm resultados excelentes.

Nesses últimos três anos, construíram e tornaram essa plataforma robusta para estar *online* e dar vazão. Há diversos centros pelo Brasil usando e testando essa plataforma, que poderia, em complementaridade com outras ferramentas de telessaúde, ajudar a levar para a ponta e resolver essas questões.

Se a gente pensar, por exemplo, no acidente vascular cerebral - os colegas médicos aqui podem dizer -, se você não tratar e identificar a causa em quatro horas, você pode ter danos irreversíveis. Você tem dois caminhos, o isquêmico e o hemorrágico, cujas formas de tratar, a princípio, podem ser antagônicas. Então, é importante você ter esse acesso ao exame, não só ao diagnóstico primário, mas ao laudo. Então, se a gente pensar, um dos principais gargalos e, talvez, o mais difícil de a gente conseguir produzir em escala não é comprar um tomógrafo, não é comprar o equipamento, é ter material humano capacitado distribuído por todo o Brasil. Se a gente quer impactar a nossa população, a gente tem que usar a tecnologia para ajudar a potencializar e gerar esse impacto.

Então, sim, está se usando, em muitas coisas, como o Rodolfo falou, e tem potencial, com muitas coisas que estão, entre outras, "pré-prontas", para fazer isso aí, escalar ainda muito mais e beneficiar a maior parte das pessoas que, hoje, dependem lá da ponta e demoram muito tempo para ver as questões resolvidas.

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

Acho que o Prof. Edson tem uma consideração também.

O SR. EDSON AMARO JÚNIOR (Para expor.) - Obrigado. Eu vou endereçar um pouco as questões anteriores, mas focar o que o Rafael Salvatore, do Amazonas, perguntou.

Rafael, a sua colocação é muito importante. Você pergunta: "Como fica a questão do direito ao livre-esclarecimento do paciente em relação à inteligência artificial?".

Esse tema é fundamental, caro e importante, mas também é desafiador. Como nós vimos aqui, é muito importante as pessoas saberem como as coisas são feitas e quais as consequências. Então, o direito ao livre esclarecimento do paciente pode ser quebrado, de novo, em vários detalhes, mas, talvez, a coisa mais importante seja: em relação à utilização da informação fornecida por inteligência artificial, o agente - o médico, o agente comunitário de saúde, a enfermeira, enfim, o ser humano que está fazendo uso dessa ferramenta - está capacitado? Ele entende essa ferramenta? Além de tudo o que nós discutimos sobre o processo de regulamentação, a pessoa da ponta consegue esclarecer isso? Claro, a nossa intenção é trazer para todo cidadão brasileiro o melhor e, se possível, melhorar o uso de recursos.

Então, agora, eu vou voltar um pouquinho, antes das duas últimas perguntas, porque o uso de IA está sendo prospectado, em muitas iniciativas, em conjunto com o Governo. A gente acabou de ver que o Fernando citou exemplos do Hospital Albert Einstein. Sim, são projetos que são feitos também dentro do SUS, mas, além de tender a verificar as possibilidades, entender os limites disso, traz-se a curva de aprendizado. Nós estamos aprendendo a fazer uso dessas ferramentas em conjunto e dentro do Sistema Único de Saúde. Acho que é um valor incomensurável para o nosso país que esse sistema tenha também uma interface com um desenvolvimento tecnológico tão avançado como esse aqui.

Mas, Rafael, a sua pergunta é pertinente, é desafiadora e não pode ser deixada de lado. Como fica a questão? A questão fica no centro das nossas conversas.

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

Vejo que o Prof. Rodolfo também quer fazer uma participação. Por favor, o senhor tem a palavra.

O SR. RODOLFO DE CARVALHO PACAGNELLA (Por videoconferência.) - Acredito que o Prof. Juliano está na frente.

O SR. JULIANO SOUZA DE ALBUQUERQUE MARANHÃO (Para expor. Por videoconferência.) - Obrigado, Rodolfo.

Queria só fazer uma consideração.

De fato, como os colegas já exemplificaram, existem várias aplicações. Agora, também existem projetos que, por vezes, não estão desenvolvidos ou enfrentam alguma hesitação em função de falta de parâmetros em relação a quais são as medidas de cuidado e as consequências, em casos de problemas ou erros, lembrando que o campo da saúde é um campo de aplicação de alto risco, um erro, um equívoco pode ter consequências graves.

Então, é muito importante que haja definição de códigos de conduta ou de parâmetros que possam servir de baliza para o desenvolvimento desses sistemas em um ambiente de segurança jurídica. Quando a gente pensa nisso, normalmente, recorre-se a uma regulação externa do Estado. Por vezes, a gente pode ter outros arranjos em que uma autorregulação é validada por uma autoridade setorial, passa a ser parâmetro e confere segurança aos atores para o desenvolvimento responsável de sistemas de inteligência artificial.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado, professor. Prof. Rodolfo.

O SR. RODOLFO DE CARVALHO PACAGNELLA (Para expor. Por videoconferência.) - Senador, quero retomar a pergunta do Sr. Andrei Luan da Silva. Ele pergunta se o SUS está preparado para receber a IA na sua realidade.

Sim, o SUS, além de estar preparado, tem um potencial muito grande para poder receber isso. Ele tem uma estrutura potencial de captação de dados, de integração de dados e um potencial de interoperabilidade de sistemas que precisa ser melhor desenvolvido, mas o SUS está ávido e pronto para isso.

Um exemplo é um trabalho que a gente está desenvolvendo em parceria com o Instituto Fernandes Figueira, da Fiocruz do Rio de Janeiro, para o Ministério da Saúde, que é um desenvolvimento de sistemas de apoio à decisão clínica para reconhecimento de condições associadas à mortalidade materna, que visa identificar situações de gravidade e reduzir as demoras associadas à morte materna. Esse tipo de sistema está em vias de ter um estudo para que seja integrado aos sistemas e-SUS que avaliam e fazem o acompanhamento pré-natal, a sistemas de monitoramento hospitalar, a sistemas de prontuário eletrônico hospitalar, para que isso possa disparar alertas.

Então, o SUS está, sim, pronto, precisa disso, está ávido e é um dos sistemas públicos, em nível mundial, mais aptos a ter esse tipo de integração.

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

Eu gostaria de colocar alguns pontos aqui para ouvir os debatedores também - quatro pontos, basicamente - sobre os quais eu fiz algumas anotações aqui.

Um é com relação à importância de esse projeto de lei não travar a inovação. Isso foi citado, acho que pelo Fernando, se não me engano, e eu concordo 100%. Essa é uma das nossas preocupações, de forma que você não tenha uma lei que seja tão detalhada, que tente cercar alguma coisa que a gente não sabe ainda. Isso é importante - isso é importante.

Se a gente verificar a curva de desenvolvimento de tecnologia, ela é basicamente uma exponencial. Se você colocar a educação na tentativa de acompanhar essa curva - esse é um dos problemas que a gente verifica, eu faço parte da Comissão de Educação aqui também -, você vê uma tentativa, ou seja, uma curva que tende a ser uma exponencial, mas que não consegue chegar ao mesmo nível, ou seja, o aluno se forma ainda com um *gap* com relação à tecnologia que ele vai encontrar no mercado de trabalho, no desenvolvimento do seu trabalho. E aí vem a regulação, a legislação, abaixo dessas duas curvas, com outro *gap* ali para acompanhar.

Eu sei que isso é desconfortável, principalmente para o pessoal mais do setor jurídico, quando a gente fala na velocidade do desenvolvimento e do conhecimento da tecnologia, porque sempre existe aquela tendência de se querer colocar todos os detalhes dentro de uma lei, mas, nesse setor, especificamente, é muito difícil você fazer esse tipo de ajuste. Ou seja, é necessário que haja outras ferramentas, de forma que a gente consiga adaptar essa legislação com o avanço, com o desenvolvimento, com as novas descobertas que vão surgir. A gente não pode travar essas novas descobertas.

Um exemplo que não tem, exatamente, a ver com isso é o marco legal das *startups*. Nós o lançamos, lá no tempo do Ministério, em torno de junho ou julho de 2021, e hoje estamos revisando esse marco aqui no Senado. Eu pedi para fazer essa revisão - já chamei algumas audiências públicas - justamente para que ele seja melhorado. E isso não vai encerrar o trabalho ali; isso precisa ser constantemente verificado e avaliado.

Eu lembro que um dos apresentadores - peço desculpas porque eu acabei não anotando o nome - fez aquela comparação com o desenvolvimento de fármacos, que, no final, você vê que tem que ser acompanhado frequentemente. É importante, realmente, que a gente continue a acompanhar essa legislação com ferramentas que a gente possa melhorar a partir do desenvolvimento que exista naquele momento. O próprio conhecimento vai levar a alterações.

Outro ponto colocado que eu achei muito importante, acho que foi pelo Fernando também, foi com relação à decisão humana.

O ser humano, em última instância, tem que ser ainda o tomador de decisão. Em grande parte as decisões que nós temos são emocionais. Elas precisam desse componente para serem justas. Às vezes, elas não são... Pode ser até que não sejam completamente lógicas, mas elas podem ser - e serão - justas por causa disso. É estranho um engenheiro falando isso, mas é importante pensar. E, quando a gente pensa na área de saúde ou na área do direito... Imagine, no direito de família, por exemplo, como se vai tomar uma decisão simplesmente com a questão lógica? Existe todo um componente emocional.

Na minha profissão, vamos supor: um piloto com 200 pessoas a bordo. Você tem um sistema de automação no avião em que ele voa sozinho, praticamente, mas, em última instância, o comandante é o responsável pelo avião. Ele é o responsável pela vida das pessoas que estão ali, e é essa pessoa que tem que tomar as decisões. O sistema pode ajudar a interpretar alguns dados, pode dar uma visão mais ampla do que está acontecendo lá fora, mas é importante que o ser humano ainda seja o protagonista de todas as decisões.

O problema que eu vejo é que nós, como seres humanos, pela natureza humana, tendemos a nos acomodar com os sistemas. Uma pergunta fácil que eu faço aqui para a pessoa imaginar o que eu estou falando é, por exemplo: quem é que usa um mapa em papel ainda para fazer navegação dentro de uma cidade, como eu via o Guia São Paulo, por exemplo? Eu praticamente não encontro mais ninguém que faça isso. Todo mundo usa o GPS no celular. Então, o sistema é muito mais cômodo, dá muito mais facilidade para você. Mas ele pode errar? Pode. Ele pode te mandar para um lugar errado, e você, seguindo aquilo lá cegamente, vamos dizer assim, sem prestar atenção ao que está acontecendo, pode ter problemas, ir aonde não deveria, e assim por diante.

Então, o ser humano tende a se acomodar. Mesmo na pilotagem de um avião, como os sistemas são automáticos, existe a tendência de os pilotos relaxarem ali, o que precisa ser compensado pelo treinamento adequado porque, se for preciso assumir o avião, em cinco segundos ou menos do que isso, você já tem que saber o que fazer naquele momento. E você é o responsável por isso.

Outro ponto que eu gostaria de colocar, que é importante também, é com relação aos centros. Nós temos diversas aplicações. Quando a gente fala de tecnologia, ela é aplicada em todos os setores. Isso é uma coisa que eu falava muito no ministério. Desde a agricultura, tecnologia na agricultura, tecnologia na defesa, na segurança pública, na educação, na indústria; todas as áreas usam tecnologia. E aqui é igual. E uma tecnologia como essa, da inteligência artificial, já está presente e vai estar cada vez mais em todas essas áreas.

Eu lembro até de que no ministério nós criamos os Centros de Inteligência Artificial. Foram oito Centros de Inteligência Artificial, com quatro áreas designadas: primeiro, era a inteligência artificial para saúde, saúde 4.0; indústria 4.0; cidades

inteligentes; e agricultura 4.0. Depois, entrou o turismo 4.0 também, como uma quinta. Eu lembro que aqueles oito centros de inteligência artificial eram cabeças de rede, e conectados com eles tinham todos os outros centros que trabalhavam naquela mesma direção.

Por que eu estou falando isso? Por causa dessa questão que foi colocada várias vezes, e eu anotei aqui, com relação à necessidade de nós termos não só... Você fala de uma agência de inteligência artificial. Mas não pode ser uma agência única de regulação em que tudo seja contido ali dentro, porque isso não vai ser possível. Pelo que eu estou entendendo, já está decantando esse tipo de informação. E a gente já tinha visto isso com relação a bancos, por exemplo, que têm o seu próprio sistema de *compliance*. Não dá para uma agência só olhar todas essas áreas ao mesmo tempo, ou seja, precisa haver uma cooperação entre os diversos órgãos que já observam cada uma dessas áreas, ou agências reguladoras; precisa-se pensar nessa governança. O mesmo problema nós enfrentávamos no Ministério de Ciência, Tecnologia e Inovações, porque o Ministério da Saúde, por exemplo, tem lá a Secretaria de Ciência e Tecnologia, o Ministério - sei lá - de Minas e Energia tinha lá a Secretaria de Ciência e Tecnologia... Como elas vão trabalhar? Elas vão trabalhar completamente descoordenadas, cada uma fazendo as coisas? Não, o ideal é que exista uma coordenação única, uma coordenação feita de forma matricial - quem trabalha com projetos sabe exatamente do que estou falando: de forma matricial -, ou seja, elas permanecem na organização estrutural e hierárquica do seu devido órgão, mas atendem a uma diretriz que vem de um órgão central, que faz essa coordenação. Acho que dessa forma a gente conseguiria ter uma coordenação mais efetiva e que atendesse as especificidades de cada uma das áreas, as quais são importantes. Na área da saúde, por exemplo, a gente está trabalhando com coisas críticas, é a saúde das pessoas.

Então, é importante que se tenha órgãos que conheçam isso e que tenham lá o seu departamento, vamos chamar assim - ou o nome que se queira colocar -, de *compliance* em inteligência artificial para aquilo que seja conectado matricialmente com essa agência geral, assim como tem a nossa agência de proteção de dados, que tem que ser conectada nesse sistema também, e essa legislação não pode ir de encontro ao que já existe em termos de legislação. Então, esse é um ponto também importante de se colocar aqui.

Eu só trouxe essas ideias todas para ouvir o comentário, porque, depois, eu vou pedir, nas considerações finais... O nosso tempo já deu aqui, mas, nas considerações finais, eu vou pedir para cada um também levar em conta isso e talvez fazer algum comentário sobre isso, dos nossos debatedores.

Eu vi ali, rapidamente, o Senador Izalci Lucas aparecendo, e depois ele sumiu. Ele está *online* ainda, com a gente? Porque eu ia pedir para ele falar. (*Pausa.*)

Já saiu?

Ah, não, ele está ali!

Antes de passar para nossos debatedores para as considerações finais, eu gostaria de passar a palavra ao Senador Izalci Lucas, que é membro da Comissão de Ciência e Tecnologia, um lutador nessa área, que está sempre trabalhando em prol da ciência e tecnologia do país. Eu tenho uma admiração muito grande desde o tempo ali como Ministro, que me ajudou, e muito, lá.

Izalci, V. Exa. tem a palavra, por favor.

O SR. IZALCI LUCAS (Bloco Parlamentar Democracia/PSDB - DF. Para discursar. *Por videoconferência.*) - Obrigado, Senador.

Bem, eu não tive oportunidade de ver a apresentação. Eu estava num debate fora do Congresso, mas vou, depois, analisar a reunião. Mas há sempre essa preocupação com relação à questão da inteligência artificial. A gente fica preocupado com o foco, com o jurídico, com quase que proibir uma coisa que a gente ainda não sabe nem o que é. Por isso, quero destacar aqui a sua preocupação, a nossa preocupação de tentar aqui sensibilizar o pessoal para deixar a coisa avançar, pelo menos para a gente saber como a gente pode liberar o setor de tecnologia e de inovação, para poder realmente o Brasil não ficar para trás nesse movimento.

Quero parabenizá-lo, e vou, depois, pedir cópia das apresentações e das notas taquigráficas para analisar. Parabéns pela promoção desta Subcomissão, que vai ter um valor imenso para nós que somos preocupados com a ciência, tecnologia, inovação e pesquisa.

Parabéns, Senador!

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado, Senador Izalci, que - novamente, reiterando aqui - é um guerreiro para a ciência e tecnologia do país. Muitas das coisas que nós temos hoje são graças ao trabalho do Senador Izalci aqui no Congresso Nacional e aqui no Senado também.

Agora, dado o nosso tempo, vou passar, então, a palavra a cada um dos debatedores para suas considerações finais. Eu peço que sejam breves nas considerações finais.

Começamos, na mesma sequência em que nós iniciamos, com o Prof. Juliano Souza de Albuquerque Maranhão, Professor da Faculdade de Direito da Universidade de São Paulo.

Professor Juliano, para suas considerações finais, por favor.

O SR. JULIANO SOUZA DE ALBUQUERQUE MARANHÃO (Para expor. *Por videoconferência.*) - Obrigado, Senador.

Eu, na verdade, vou fazer coro a uma fala colocada pelo senhor na abertura da Comissão na terça-feira, porque, para mim, foi significativa a ênfase nos benefícios da tecnologia antes da preocupação com os riscos, porque um dos riscos em relação à inteligência artificial está na subutilização, em não aproveitar os benefícios potenciais que essa tecnologia tem a trazer e que já vem produzindo. O campo da medicina é um exemplo, mas isso ocorre em vários campos da atividade econômica e da atividade social.

Dessa forma, pensar também na regulação como um veículo para fomento ao investimento e fomento ao desenvolvimento tecnológico do país para inserção na economia global também é importante, bem como pensar em formas flexíveis para que haja o desenvolvimento de uma inteligência artificial responsável no país, mas também com a flexibilidade para adaptação contínua da regulação ao seu estado da arte.

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

Passo a palavra agora ao Prof. Edson Amaro Júnior, Professor Livre-Docente da Faculdade de Medicina da Universidade de São Paulo.

Professor, para suas considerações.

O SR. EDSON AMARO JÚNIOR (Para expor.) - Acho que, exatamente seguindo aí o raciocínio já colocado, o esforço de regulação deve se pautar para que nós tenhamos normas principiológicas, porque é importante nós termos a visão de que os processos, as atitudes e a forma com que a gente utiliza a tecnologia sejam, obviamente, feitos de maneira responsável pela sociedade como um todo; entretanto, com o cuidado de que a gestão não só do desenvolvimento e da implementação, mas, principalmente, da monitorização e fiscalização seja feita de maneira matricial, como já foi colocado também.

Reforço a ênfase de que essa tecnologia é necessária para a prática da boa medicina, do que nós temos hoje no conhecimento da humanidade. Isso não é algo que a gente tem o luxo de deixar de utilizar. Por outro lado, a forma com a qual nós podemos fazer, de novo, o melhor uso e obter os benefícios tão claramente apontados pelas pesquisas atuais envolve, sim, uma articulação um pouco mais ampla e menos centralizada do que a gente tem observado aqui.

Mais uma vez, agradeço e parablenizo a iniciativa.

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

Passo a palavra, para suas considerações finais, ao Prof. Rodolfo de Carvalho, Professor da Universidade Estadual de Campinas.

Professor, tem a palavra para as considerações finais.

O SR. RODOLFO DE CARVALHO PACAGNELLA (Para expor. *Por videoconferência.*) - Senador, novamente, agradeço a possibilidade de estar aqui neste debate.

Eu vou retomar uma anedota que um dos artigos que eu trouxe e que embasaram minha fala traz, que é uma verdade inconveniente sobre a utilização de IA na saúde.

No século XIX, a Inglaterra fez um investimento vultoso em saneamento nas margens do rio em Londres, e esse investimento era muito questionado naquele momento. Hoje, ninguém em sã consciência questiona um investimento importante, grande, em saneamento nas cidades.

O investimento que a gente deve fazer na área de inteligência artificial para a saúde hoje é comparável. A gente precisa fazer um investimento grande hoje, porque, no futuro, a gente vai reconhecer que isso é fundamental; que hoje a gente não tem como mais pensar saúde sem o apoio de algumas dessas ferramentas.

Portanto, precisamos investir energia, envolvimento, regulação, para que isso funcione de maneira adequada.

Muito obrigado.

O SR. PRESIDENTE (Astronauta Marcos Pontes. Bloco Parlamentar Vanguarda/PL - SP) - Obrigado, Prof. Rodolfo de Carvalho.

Eu passo a palavra agora para o Prof. Wagner Meira Júnior, Professor Titular do Departamento de Ciência da Computação da Universidade Federal de Minas Gerais.

Professor, suas considerações finais.

O SR. WAGNER MEIRA JÚNIOR (Para expor. *Por videoconferência.*) - Obrigado, Senador.

Eu esqueci de mencionar que eu sou pesquisador principal de um centro desses de IA para a saúde.

Uma pergunta que eu achei bem pertinente é a questão da inovação e da incorporação da decisão humana. Eu acho que essas duas coisas caminham juntas quando nós falamos de IA em saúde. Por quê? Porque, se um profissional de saúde não incorporar, na sua prática, a utilização daquela tecnologia - e para isso há uma demanda brutal de requalificação, a gente já tem percebido isso aqui no nosso centro, ou seja, metade dos nossos pesquisadores são da área de saúde -, nós não vamos realmente conseguir inovar. Nós vamos continuar naquela lógica de que teremos soluções prototípicas tecnológicas bastante arrojadas, mas que, na prática, não fazem a diferença.

Para o setor de saúde, o humano no *loop* é o profissional de saúde e, aproveitando a participação do senhor na Comissão de Educação, eu acho que é importante que isso seja levado lá.

Ou seja, nós precisamos de uma mudança fundamental na perspectiva que as várias áreas, em particular as áreas de risco, têm sobre a inteligência artificial.

Muito obrigado.

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

Realmente, essa é uma das discussões que a gente tem tido. Aliás, eu até criei uma frente aqui no Congresso, a Frente Parlamentar em Favor da Educação Profissional e Tecnológica.

A ideia toda é para a gente ver esse *gap* que nós temos muito grande na formação de pessoas, principalmente para a área de tecnologia de informação, comunicação etc. - e, por outro lado, quanta gente precisando de emprego no país. Ter esse meio de campo preenchido aqui para que essas pessoas que precisam de emprego sejam qualificadas para atender o mercado de trabalho com as tecnologias emergentes é extremamente importante. É uma consideração que a gente não pode colocar para trás.

Isso é urgente e agora, como a gente tem colocado aqui na Comissão.

Então, obrigado pela consideração.

Eu passo agora a palavra para o Prof. Fernando José Ribeiro Sales, Professor do Departamento de Engenharia Biomédica da Universidade Federal de Pernambuco, para suas considerações finais.

O SR. FERNANDO JOSÉ RIBEIRO SALES (Para expor. *Por videoconferência.*) - Obrigado, Senador.

Eu queria deixar aqui uma reflexão para os senhores, ilustríssimos e ilustríssimas Parlamentares: pensarem em qual legado vocês querem deixar com esse projeto de lei, com essa proposta. A gente sabe que o processo legislativo é longo, mas eu acho que, se eu estivesse aí na posição de vocês, eu focaria talvez, neste momento, nessa questão dos princípios de fato, para garantir essa questão de responsabilizar, identificar os limites mais máximos à fronteira e criar ambientes, estruturas como os *sandboxes*, para que a gente possa experimentar e ir aprendendo a partir dessa forma. O método científico é isto: a gente cria hipóteses, define um método, avalia o resultado e vai sempre realimentando.

Então, tomar o cuidado só de, eventualmente, na proposta de projeto de lei, com tudo o que lá está, a gente não ter o efeito contrário de acabar travando. Esse é um assunto tão sensível e estratégico para o nosso Brasil, essa é uma pauta de Estado, essa é uma pauta que vai alimentar e que vai alavancar muito do que a gente tem de potencial em várias áreas, e esse olhar, talvez, quando a gente não sabe direito o que fazer, se regulamentar dizendo talvez seja pior do que deixar ali e prever.

Então, acho que é importante focar nesses princípios, mas eu gostaria de parabenizá-los pela nobre iniciativa de ter esse processo de escuta amplo para vários setores, acho que isso é muito rico dentro do processo legislativo e desejo sabedoria nos próximos passos.

Muito obrigado.

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

Passo agora a palavra para o Tiago José de Carvalho, Diretor de Tecnologia da empresa Inovia.

Tiago, para as suas considerações finais.

Obrigado.

O SR. TIAGO JOSÉ DE CARVALHO (Para expor. *Por videoconferência.*) - Obrigado, Senador.

Queria agradecer novamente a oportunidade de estar aqui discutindo com os colegas. Isso é muito rico, isso agrega muito à democracia, traz muito ganho para toda a construção do processo. E a minha mensagem final vai em duas linhas. Os colegas já tocaram no ponto, mas mais para reforçar.

A primeira é que a minha visão de qualquer sistema de computação, informática e a inteligência artificial está ali dentro, é que a gente sempre está no meio, e não na ponta. Então, a IA vai servir como ferramenta para o tomador de decisão, para melhorar, para ajudar o tomador de decisão. Então, o projeto de lei acredito que deveria refletir isso, ele deveria se preocupar com isso.

E o segundo ponto, pegando a fala do Senador que disse sobre a matriz de instrumento do tomador de decisão, é que eu gostaria que o projeto de lei olhasse para outras coisas que já foram feitas e que visam a fomentar a inovação no país, como, por exemplo, o marco legal das *startups*. Eles têm que conversar, o projeto de lei que regulamentar a inteligência artificial tem que conversar com o marco legal das *startups*, senão vão ser duas coisas, uma brigando com a outra, o que vai travar a inovação, o que vai travar o progresso.

Então, é isso, essas são as minhas considerações.

Obrigado pela oportunidade novamente.

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

A gente estava comentando justamente isso que eu falei com o Prof. Edson aqui: que um dos meus trabalhos aqui no Senado tem sido justamente este, de trazer essas diversas legislações diferentes para que sejam discutidas ou rediscutidas aqui e para que elas sejam aperfeiçoadas, mas também alinhadas.

Por exemplo, o marco legal das *startups* junto com a Lei do Bem. Nós estamos agora nessas discussões. Eu fui o Relator da Lei do Bem, o autor da modificação, o Senador Izalci, estava aqui conosco recentemente. E, depois, o marco legal da inovação. Tem a Lei de TICs, ou a lei da informática, e todas essas leis, essas legislações precisam ser alinhadas.

Hoje mesmo, de manhã, aqui, tratamos também de inteligência artificial - nós tivemos representantes da Agência Nacional de Proteção de Dados - e da necessidade que existe do alinhamento dessa legislação, de forma que você não repita o que já está escrito em outra legislação ou, pior ainda, que você não contradiga uma coisa e faça uma que vá de encontro à outra, o que seria muito ruim em termos de legislação. Então, é muito importante que nós tenhamos esses alinhamentos.

Parabéns!

Eu gostaria aqui de, finalmente, agradecer a presença de todos os nossos debatedores, todos os nossos apresentadores, agradecer a todos aqueles que nos acompanharam através da TV Senado, através das redes, aqueles que mandaram suas perguntas através do e-Cidadania ou pelo telefone e dizer que nós estamos aqui à disposição para que nós tenhamos o melhor em termos de legislação através da participação da sociedade civil.

É muito importante ter esses debates, essas conversas e colocarem-se os pontos. Dessa forma, a gente consegue o melhor em termos de legislação, que certamente não vai ser perfeita, como nada é, mas precisamos ter em mente que ela precisa ser, ao longo do tempo, melhorada, aperfeiçoada, e esse trabalho tem que ser contínuo.

Obrigado a todos.

Não havendo mais nada a tratar nesta reunião, eu a declaro encerrada, com o desejo de um ótimo final de semana para todos.

Obrigado.

(Iniciada às 14 horas e 27 minutos, a reunião é encerrada às 16 horas e 17 minutos.)