



PONTIFÍCIA UNIVERSIDADE CATÓLICA DE GOIÁS
PRÓ-REITORIA DE GRADUAÇÃO
ESCOLA DE DIREITO, NEGÓCIOS E COMUNICAÇÃO
CURSO DE DIREITO
NÚCLEO DE PRÁTICA JURÍDICA
COORDENAÇÃO ADJUNTA DE TRABALHO DE CURSO I

A INTELIGENCIA ARTIFICIAL APLICADA ÀS RELAÇÕES EMPRESARIAIS
ANÁLISE NORMATIVA

ORIENTANDO: DOUGLAS GABRIEL MARQUES NERES
ORIENTADOR: JOÃO BATISTA VALVERVE

Goiânia/GO

2024

DOUGLAS GABRIEL MARQUES NERES

**A INTELIGENCIA ARTIFICIAL APLICADA ÀS RELAÇÕES EMPRESARIAIS
ANÁLISE NORMATIVA**

Artigo científico apresentado à disciplina Trabalho de Curso I, da Escola de Direito, Negócios e Comunicação da Pontifícia Universidade Católica de Goiás (PUCGOIÁS). Prof. Orientador (a) - João Batista Valverde de Oliveira.

Goiânia/GO

2024

RESUMO

A inteligência artificial já está consolidada em todos os territórios mundiais, alguns permitem mais ou menos investimentos para uso do próprio Estado e da sociedade em geral. O conceito de IA é amplo, não se trata de uma ferramenta ou apenas um software, embora não seja um atributo meramente determinável, trata-se de um complexo de “um pouco de tudo”, é uma compilação do ser humano em forma de robô, seja em forma material (físico) ou imaterial (software) e até mesmo os dois em conjunto. O que vale destacar é que a ideia por trás é fazer um ser humano em forma de máquina, e este possui competência para aprender e agir sem limitações, incansável, imparável.

Com efeito, por mais que a IA vem crescendo significativamente e já esteja consolidada mundialmente, a sua própria evolução tem impactado a sociedade do consumo de bens e serviços, o crescimento da IA é ilimitado, contudo, necessário se faz regulamentar o uso deste complexo autônomo-inteligente para que, cedo ou tarde, não se torna incontrolável e cause danos irreparáveis à sociedade *latu sensu*.

A dificuldade do legislador pátrio é tamanha que, ao analisar brevemente o contexto histórico do marco da IA na sociedade, percebe-se que, por exemplo, a primeira tentativa de regular a IA no Brasil foi em meados do ano de 2020, por um Deputado Federal, o que, inclusive o referido projeto de lei em epígrafe já fora rejeitado por maioria absoluta da Câmara dos Deputados. Pouco tempo depois, em 2021, um Senador Federal tomou a iniciativa de propor um projeto de lei para regular a IA. No ano de 2023, percebendo, pois, um outro Senador, uma chance de melhorar o projeto daquele primeiro, formulou um outro projeto e, ambos os projetos destes senadores, hodiernamente, no segundo trimestre de 2024, estão aguardando parecer de uma comissão interna, cuja qualificação dos membros ditos especialistas deste complexo-inteligente é notória.

Assim, temos que essas várias tentativas de regular o uso da IA no território brasileiro, por parte do legislador, é de extremada dificuldade, uma vez que, na medida em que o legislador regula os atributos normativos da IA no ano de 2023, por exemplo; em 2024 ou em 2025, a IA já avançou de forma estupenda, a ponto de que o ser humano chamado legislador não consiga acompanhar e regula-la de forma eficiente.

Será que em breve, precisaremos da IA para criar normas para a IA?

SUMÁRIO

INTRODUÇÃO	4
1 - DIRETRIZES PARA O USO ÉTICO DA INTELIGÊNCIA ARTIFICIAL	5
2 - BASES DE UMA IA DE CONFIANÇA.....	5
2.1 Princípios éticos no contexto dos sistemas de IA.....	5
2.2 Respeito a autonomia humana	5
2.3 Prevenção de danos	6
2.4 Equidade.....	6
2.5 Explicabilidade	6
3 - CONCRETIZAÇÃO DE UMA IA DE CONFIANÇA	7
3.1 Intervenção Humana	9
3.2 Fiscalização Humana	9
3.3 Controle	10
4 - AVALIAÇÃO DE UMA INTELIGÊNCIA ARTIFICIAL DE CONFIANÇA	11
5 - LEGISLAÇÃO APLICÁVEL.....	12
CONCLUSÃO.....	13
REFERÊNCIAS	14

INTRODUÇÃO

Hodiernamente os atributos da IA (inteligência artificial) são vistos em diversos ambientes empresariais e domésticos. Em ambientes residenciais, é comum ter casas inteligentes, ou seja, aquelas que identificam vozes para realizar uma tarefa como acender luzes, acionam os robôs para providenciarem a limpeza ao identificarem um corpo estranho, regulam a temperatura do ambiente conforme programação etc.

Já nos ambientes corporativos, observamos a aplicação da IA no marketing, na gestão financeira, nos Recursos Humanos, na segurança e na tomada de decisão sobre produtos e mercadorias de investimentos.

Buscando elaborar certos princípios de governança digital e corporativa, aplicados à Inteligência Artificial nas relações empresariais e servindo como orientações não obrigatórias para os legisladores elaborarem as suas leis, em junho de 2018 a Comissão Europeia selecionou peritos (GPAN IA) de notório saber tecnológico de informação e reputação ilibada para discutirem os parâmetros norteadores aplicáveis a IA. Criou-se então a *Ethics Guidelines For Trustworthy AI* em abril de 2019, que em uma tradução livre significa: Diretrizes de ética para uma Inteligência Artificial confiável.

Essas diretrizes serão a base de todo o arcabouço normativo, em respeito a aplicação da Inteligência Artificial nos produtos e serviços das atividades empresárias sem perder a característica da sua função social, conforme dispõe a matéria do site X2 inteligência digital. (Históriadainteligênciaartificial.X2inteligênciadigital, 2022. Disponível em: <https://x2inteligencia.digital/2020/02/20/historia-da-inteligenciaartificial-2/>. Acesso em 15/04/2022).

A INTELIGENCIA ARTIFICIAL APLICADA ÀS RELAÇÕES EMPRESARIAIS

ANÁLISE NORMATIVA

1 - DIRETRIZES PARA O USO ÉTICO DA INTELIGÊNCIA ARTIFICIAL

De acordo com as Diretrizes da *Ethics Guidelines For Trustworthy AI*, a Inteligência Artificial confiável deve ser:

- 1) legal - respeitando todas as leis e regulamentos aplicáveis;
- 2) ético - respeitando os princípios e valores éticos, por mais que a legislação possa estar defasada;
- 3) sólida - tanto do ponto de vista técnico quanto levando em consideração seu ambiente social.

Os três elementos são relevantes para, ao menos, orientar o aplicador da inteligência artificial, (p. ex., o âmbito e o conteúdo da legislação em vigor podem não se coadunar, por vezes, com as normas éticas). Temos a responsabilidade individual e coletiva, enquanto sociedade, de procurar garantir que os três componentes contribuam para assegurar uma IA de confiança.

Para uma melhor compreensão das diretrizes da **EGFT – AI**, fora criado um relatório que objetivou fornecer um norte básico, composto por peritos especialistas que abordaram 3 (três) fontes, sendo-as analisadas de forma sucinta e eficiente os seguintes: Bases, Concretização e avaliação de uma IA de confiança.

2 BASES DE UMA IA DE CONFIANÇA

Uma observação a ser feita é que os consumidores da IA, só a vão consumi-la se houver clareza de como ela funcionará, pelo exposto, constata-se que o termo base se refere às principais estruturas que sustentam o uso e funcionamento da aplicação da inteligência artificial, *in verbis*.

2.1 PRINCÍPIOS ÉTICOS NO CONTEXTO DOS SISTEMAS DE IA

Enraizados nos direitos fundamentais, os princípios a seguir devem ser respeitados para assegurar que os sistemas de IA sejam desenvolvidos, implantados e utilizados de forma confiável. São especificados como imperativos éticos, que os profissionais no domínio da IA devem esforçar-se sempre por respeitar. Segundo a *Ethics Guidelines For Trustworthy AI*, estes são os princípios:

2.2 RESPEITO A AUTONOMIA HUMANA

Os sistemas de IA não devem subordinar, coagir, enganar, manipular, condicionar ou arregimentar injustificadamente os seres humanos. Em vez disso, os sistemas de IA devem ser concebidos para aumentar, complementar e capacitar as competências cognitivas, sociais e culturais dos seres humanos.

2.3 Respeito a autonomia humana Prevenção de danos

Os sistemas de IA não podem afetar negativamente os seres humanos de qualquer forma. Isto implica a proteção da dignidade, bem como da integridade mental e física, do ser humano. Os sistemas de IA e os ambientes em que operam devem ser seguros e protegidos.

2.4 Equidade

O princípio da equidade sobre as dimensões substantiva e processual. A substantiva implica um compromisso com a garantia de uma distribuição equitativa e justa dos benefícios e dos custos, bem como de inexistência de enviesamentos injustos, discriminação e estigmatização contra pessoas e grupos. A dimensão processual da equidade implica uma possibilidade de contestar e procurar vias de recurso eficazes contra as decisões tomadas por sistemas de IA e pelos seres

humanos que os utilizam. Para tal efeito, a entidade responsável pela decisão deve ser identificável e os processos decisórios explicáveis.

2.5 Explicabilidade

A explicabilidade é crucial para criar e manter a confiança dos utilizadores dos sistemas de IA. Logo, os processos de utilização para uma determinada tarefa e conclusão desta têm de serem transparentes, explicando aos usuários o percurso que o sistema vai utilizar até a consecução do seu fim.

Estes casos são designados por algoritmos de caixa negra e exigem especial atenção. Nessas circunstâncias, podem ser necessárias outras medidas da explicabilidade (p. ex., a rastreabilidade, a auditabilidade e a comunicação transparente sobre as capacidades do sistema), desde que o sistema, no seu conjunto, respeite os direitos fundamentais.

Por exemplo, as recomendações de compra inexatas geradas por um sistema de IA podem suscitar poucas preocupações éticas, ao contrário dos sistemas de IA que avaliam se uma pessoa condenada por uma infração penal deve ou não ser beneficiada pela liberdade condicional.

3 CONCRETIZAÇÃO DE UMA IA DE CONFIANÇA

O termo concretização se torna plausível a partir do momento em que os princípios deixam a teoria e se efetiva no mundo das partes interessadas nos sistemas de IA, de acordo com *Ethics Guidelines For Trustworthy AI*, as partes interessadas são todas as pessoas envolvidas com a inteligência artificial, ou seja, aqueles que a criaram por intermédio de software ou hardware, aqueles que a implantaram nos processos empresariais, bem como os consumidores *standard*, ou seja, os consumidores finais da tecnologia.

Para concretizar os princípios em normas garantidoras, os responsáveis pela criação da IA devem produzir e programar de acordo com os princípios basilares; aos que aderiram ao processo tecnológico devem garantir, veementemente, que houve a concretização e obediência das normas; já os consumidores *standard* devem ser informados e exigir que os preceitos normativos estão de acordo e que os seus direitos estão garantidos.

De acordo com a *Ethics Guidelines For Trustworthy AI* a lista de requisitos a seguir apresentada não é exaustiva. Inclui aspetos sistêmicos, individuais e sociais:

- 1-Ação e supervisão humanas: corresponde à inclusão dos direitos fundamentais, a ação humana e a supervisão humana;
- 2-Solidez técnica e segurança: incluindo a resiliência perante ataques e a segurança, os planos de recurso e a segurança geral, a exatidão, a fiabilidade e a reprodutibilidade;
- 3-Privacidade e governação dos dados: incluindo o respeito da privacidade, a qualidade e a integridade dos dados e o acesso aos dados;
- 4-Transparência: Incluindo a rastreabilidade, a explicabilidade e a comunicação;
- 5-Diversidade, não discriminação e equidade: incluindo a prevenção de enviesamentos injustos, a acessibilidade e a concessão universal e a participação das partes interessadas;
- 6-Bem-estar societal e ambiental: incluindo a sustentabilidade e o respeito do ambiente, o impacto social, a sociedade e a democracia;
- 7-Responsabilização: incluindo a auditabilidade, a minimização e a comunicação dos impactos negativos, as soluções de compromisso e as vias de recurso.

Ocorre que, com amparo na *Ethics Guidelines For Trustworthy AI* e na legislação pertinente ao tema (13.709/18 – LGPD e os projetos de lei que estão em análise no Senado Federal), Faz-se necessário se socorrer a outras áreas do conhecimento de boas práticas empresariais, oportuno ressaltar a Governança Corporativa, senão vejamos a orientação que o professor de administração, economia

e contabilidade da Universidade de São Paulo (FEA/USP) Alexandre Di Miceli traz sobre a Governança Corporativa.

O que chamamos de “governança corporativa” diz respeito a um conjunto de práticas de negócio alicerçadas sobre princípios comuns que foram desenvolvidas em todo o mundo a partir do início da década de 1990. A discussão sobre o tema se inicia, portanto, a partir dos sólidos princípios que desde o início nortearam o movimento em prol da boa governança. (SILVEIRA, 2014, p. 02).

Pois bem, inseridos nas práticas de governança corporativa, encontram-se métodos éticos de informação, trata-se da governança digital. Mestre no assunto é o filósofo italiano Luciano Floridi, na qual estabeleceu alguns parâmetros para auxiliar no desenvolvimento saudável da Inteligência Artificial com os seres humanos que a utilizam. Vejamos então, o conceito de governança digital, segundo Floridi:

Governança digital é a prática de estabelecer e implementar políticas, procedimentos e padrões para o desenvolvimento, uso e gestão adequados da infosfera (complexo ambiente informacional). É também uma questão de convenção e boa coordenação, às vezes nem moral nem imoral, nem legal nem ilegal. Por exemplo, por meio da governança digital, uma agência governamental ou uma empresa pode (i) determinar e controlar processos e métodos usados por administradores de dados e custodiantes de dados para melhorar a qualidade, confiabilidade, acesso, segurança e disponibilidade de seus serviços; e (ii) elaborar procedimentos eficazes para a tomada de decisões e para a identificação de responsabilidades com relação aos processos relacionados a dados. (FLORIDI, 2018, p. ?)

A governança digital possibilita às partes interessadas que, primeiramente, ou seja, antes da aplicação da inteligência artificial, seja elaboradas as diligências necessárias a fim de garantir que nenhum erro venha a causar danos irreparáveis aos usuários, quando em uso constante, nas relações com os seres humanos que desejem algum serviço. É oportuno ressaltar que, após a configuração das práticas de governança digital será necessário conferir se não há nenhum empecilho que prejudique a governança corporativa da empresa, pois, se não prejudicar, significa que o objeto social da empresa está de acordo com a sua função social e, necessariamente não haverá maiores problemas no ciclo de atendimento da IA com os usuários.

E no ato de configurar os sistemas inteligentes para trabalhar, a *Ethics Guidelines For Trustworthy AI* nos orienta sobre os métodos que as ciências da computação utilizam a fim de garantir a integralidade das partes interessadas sempre

levando em consideração o princípio do respeito da autonomia humana, na ocasião, a interferência humana se torna necessária, observe-se.

A supervisão pode ser realizada mediante mecanismos de governança como as abordagens de intervenção humana (*human in the loop*), de fiscalização humana (*human on the loop*), ou de controle humano (*human in command*).

3.1 INTERVENÇÃO HUMANA

A abordagem HITL refere-se à capacidade de intervenção humana em todos os ciclos de decisão do sistema, a qual, em muitos casos, não é possível nem desejável, sobretudo porque se desvincularia da lógica da IA.

Com efeito, o Projeto de Lei nº 2.338, de 2023, de iniciativa do Senador Rodrigo Pacheco (PSD/MG), dispõe na seção III, com ênfase no seu art. 9º, sobre a possibilidade de o usuário solicitar a intervenção humana, senão vejamos:

Seção III Do direito de contestar decisões e de solicitar intervenção humana
Art. 9º A pessoa afetada por sistema de inteligência artificial terá o direito de contestar e de solicitar a revisão de decisões, recomendações ou previsões geradas por tal sistema que produzam efeitos jurídicos relevantes ou que impactem de maneira significativa seus interesses.

Não obstante, sobre a faculdade da intervenção humana é preciso não perder de vista que, por um lado, deixa de existir a autonomia do sistema inteligente e, por outro, faz-se necessário tendo em vista que, ao menos nos primeiros momentos de aplicação, esse atributo negativo garante ao usuário consumidor uma segurança de que ao usar o sistema, tenha a cognição de que os riscos e possíveis prejuízos inerentes aos erros da máquina podem ser evitados ou amenizados com a fiscalização de um expert treinado para suprir as necessidades que a máquina não consegue ou que ainda não aprendera a executar.

3.2 FISCALIZAÇÃO HUMANA

A abordagem HOTL refere-se à capacidade de intervenção humana durante o ciclo de concessão do sistema e de acompanhamento do funcionamento do sistema.

Da mesma maneira, é preciso deixar a faculdade aos usuários de solicitarem a intervenção humana sempre que há certa insegurança no usufruto ou até mesmo quando o sistema apresentar instabilidade que o torne ineficaz.

3.3 CONTROLE HUMANO

A abordagem HIC refere-se à capacidade de supervisionar toda a atividade do sistema de IA (incluindo o seu impacto económico, societal, jurídico e ético mais geral) e de decidir quando e como utilizar o sistema em qualquer situação específica. Tal pode incluir a decisão de não utilizar um sistema de IA numa determinada situação, de estabelecer níveis de apreciação humana durante a utilização do sistema, ou de assegurar a capacidade de anular uma decisão tomada por um sistema. Além disso, deve ser garantido que as autoridades públicas responsáveis pela aplicação da lei têm a possibilidade de exercer a supervisão em conformidade com o seu mandato. Podem ser necessários mecanismos de supervisão em graus variáveis para apoiar outras medidas de segurança e controle, dependendo do domínio de aplicação e do potencial risco do sistema de IA. Não havendo alteração das demais condições, quanto menor for a supervisão que um ser humano pode exercer sobre um sistema de IA, maior será a necessidade de sujeitar o mesmo a amplos testes e a uma governação rigorosa.

A intervenção e a fiscalização humana, (*human in the loop* e *human on the loop*), respectivamente, não faz o menor sentido em um sistema de inteligência artificial, haja vista, como já ressaltado, desvincularia com o objetivo de trazer autonomia à máquina, para que ela possa realizar atividades como humanos, assim como previa *Alan Turing*, no sentido de tornar a máquina, um “humano”, e poder conversar com o sistema como se ele tivesse vida própria.

Por fim, resta-nos a intervenção por intermédio do controle humano (*human in command*), a capacidade de supervisionar e decidir quando e como utilizar o sistema de inteligência é benéfica à empresa e aos usuários que o utilizam, pois possibilitam a autonomia da IA e o julgamento pessoal do controlador sobre as questões relativas a melhor oportunidade de usufruir do potencial da tecnologia nas relações comerciais. O projeto de lei do Senado Federal número 5.051 de 2019 adotou este modo de supervisão humana e, conseqüentemente, a responsabilidade objetiva pelos danos causados pela tecnologia, conforme observação do art. 4º do projeto de

lei em trâmite. O que, apesar de ser o método ideal, causa certo constrangimento desnecessário, haja vista que, a interferência das autoridades governamentais que elaboram a lei poderá supervisionar a aplicação da IA.

Tal interferência fere a livre concorrência do mercado e a autonomia empresarial. Vale mencionar a orientação feita por Ludwig Von Mises na qual “[...] a sociedade aberta é impossível sem a lógica competitiva. Sem mercado não existe sociedade aberta. O ressentimento contra o mercado é o ressentimento contra a humanidade.” (MISES *apud* MAZZILLI, 2010, p.34).

4 AVALIAÇÃO DE UMA INTELIGÊNCIA ARTIFICIAL DE CONFIANÇA

Por fim, estabelece a *Ethics Guidelines For Trustworthy AI* que, a cooperação entre os setores da empresa deve ser observada para agirem como um corpo gestor em plena e constante harmonia, a seguir, verifica-se, em um rol não exaustivo, a competência de cada departamento interno da empresa.

Conselho de Administração, neste caso, a gestão de topo debate e avalia o desenvolvimento, a implantação ou a aquisição de IA e constitui a instância superior para avaliar todas as inovações e utilizações de IA, quando são detectadas preocupações críticas. Envolve as pessoas afetadas pela eventual introdução de sistemas de IA (p. ex., os trabalhadores) e os seus representantes ao longo de todo o processo, através de procedimentos de informação, consulta e participação.

Departamento jurídico, responsável por controlar a utilização da lista de avaliação e a sua necessária evolução para acompanhar as mudanças tecnológicas e regulamentares. Atualiza as normas ou políticas internas relativas aos sistemas de IA e garante que a utilização de tais sistemas está conforme com o quadro legal e regulamentar em vigor e com os valores da organização.

Desenvolvimento de produtos, trata-se de um departamento de desenvolvimento de produtos e serviços utiliza a lista de avaliação para avaliar os produtos e serviços baseados em IA e regista todos os resultados. Estes resultados são debatidos a nível da gestão, que aprova em última instância as aplicações baseadas em IA, novas ou revistas.

Garantia da Qualidade, um atributo que garanti a qualidade (ou equivalente) verificando os resultados da lista de avaliação e toma medidas para levar uma questão ao nível de gestão superior, se o resultado não for satisfatório ou se forem detectados resultados imprevistos.

Recursos Humanos, com ênfase no departamento de recursos humanos, assegura a combinação apropriada de competências e a diversidade de perfis dos criadores de sistemas de IA. Assegura que é ministrado o nível de formação adequado sobre a IA de confiança dentro da organização.

Operações correntes, aqui os criadores e gestores de projetos incluem a lista de avaliação no seu trabalho quotidiano e documentam os resultados e as conclusões da avaliação. Ao utilizar a lista de avaliação na prática, a *Ethics Guidelines For Trustworthy AI* recomenda que se preste atenção não só aos domínios que suscitam preocupação, mas também às perguntas que não podem ser (facilmente) respondidas.

Um potencial problema poderá ser a falta de diversidade das capacidades e competências da equipe que está a desenvolver e a testar o sistema de Inteligência Artificial, podendo ser necessário envolver outras partes interessadas internas ou externas à organização. Recomenda-se vivamente o registo de todos os resultados, tanto em termos técnicos como em termos de gestão, assegurando que a resolução de problemas pode ser compreendida a todos os níveis da estrutura de governação.

5 LEGISLAÇÃO APLICÁVEL

Hodiernamente, no Brasil, há alguns projetos de leis que objetivam regular o uso da inteligência artificial no território nacional, o primeiro deles, mas já rejeitado por unanimidade pela mesa diretora da Câmara dos Deputados, foi o Projeto de lei nº 21/2020 de autoria do Deputado Federal Eduardo Bismarck - PDT/CE, cuja Ementa Estabelecia o conjunto de princípios, direitos e deveres para o uso de inteligência artificial no Brasil, e dava outras providências, conforme informativo do site da Câmara dos Deputados.

No âmbito do Senado Federal, temos o projeto de Lei nº 872, de 2021. Cujas iniciativa foi do Senador Veneziano Vital do Rêgo (MDB/PB) a Ementa: Dispõe sobre o uso da Inteligência Artificial. Explicação da Ementa: Dispõe sobre os marcos éticos e as diretrizes que fundamentam o desenvolvimento e o uso da Inteligência Artificial no Brasil.

Apesar de ser o primeiro projeto de lei do Senado, tendente a normatizar o uso da IA, tal projeto de lei está majoritariamente negativado pela cidadania nacional, contudo, no dia 21/02/2024 fora encaminhado a CTIA do Senado (Comissão Temporária Interna sobre Inteligência Artificial no Brasil) a fim de elaborar parecer sobre a matéria do projeto, e, na data de 08/04/2024, ainda está aguardando o respectivo documento.

Com efeito, o Projeto de Lei nº 2.338, de 2023, de iniciativa Senador Rodrigo Pacheco (PSD/MG) Ementa: Dispõe sobre o uso da Inteligência Artificial, é o projeto de lei que mais chamou a atenção do cidadão conforme votação geral disposta no site do Senado Federal no dia 08/04/2024 com posição favorável ao projeto de mais de 220% com relação aos que não sustentam ser o referido projeto, o mais adequado.

Contudo, da mesma forma que o projeto mencionado anteriormente, no dia 21/02/2024 fora encaminhado CTIA do Senado (Comissão Temporária Interna sobre Inteligência Artificial no Brasil) a fim de elaborar parecer sobre a matéria do projeto, e, na data de 08/04/2024, ainda está aguardando o respectivo documento.

CONCLUSÃO

De acordo com as Diretrizes da *Ethics Guidelines For Trustworthy AI*, a Inteligência Artificial confiável deve ter os atributos da legalidade, eticidade e solidez, além de outros nos quais visam a tutela dos interesses dos usuários, ao menos enquanto o sistema ainda deixar uma margem de ceticismo, para tanto, as diretrizes gerais europeia fez uma abordagem significativa, apresentando sugestão da interferência humana antes, no curso e depois da tomada de decisão da máquina.

Assim, a supervisão, neste caso, pode ser realizada mediante mecanismos de governança como as abordagens de intervenção humana (*human in the loop*), de fiscalização humana (*human on the loop*), ou de controle humano (*human in command*), o que gera confiança e credibilidade ao usuário consumidor enquanto a legislação pátria não define às normas de regulação e fiscalização da inteligência artificial no Brasil.

REFERÊNCIAS

FLORIDI L. Soft ethics, the governance of the digital and the General Data Protection Regulation. Phil. Trans. R. Soc. 2018, A 376: 20180081. Disponível em: [http://dx.doi.org/10.1098/rsta.2018.0081]. Acesso em: 02/07/2022.

Hobbes, Thomas, 1588-1679. *Leviatã, ou, matéria, forma e poder de um Estado eclesiástico e civil* / Thomas Hobbes; tradução Rosina D'Angina ; consultor jurídico Thélío de Magalhães. – São Paulo: Martin Claret, 2009. – (Coleção a obra prima de cada autor. Série ouro; 1).

Lee, Kai-Fu. **Inteligência artificial**: como os robôs estão mudando o mundo, a forma como amamos, nos relacionamos, trabalhamos e vivemos - 1. ed. - Rio de Janeiro : Globo Livros, 2019.

MAZZILLI, Marcelo. **Estado? não obrigado**. 1ª edição. São Paulo: Instituto Ludwig von Mises Brasil, 2010.

[PL 2338/2023 - Senado Federal](#) acesso em 08/04/2024 às 21:24.

[PL 872/2021 - Senado Federal](#) acesso em 08/04/2024 às 20:57.

[Portal da Câmara dos Deputados \(camara.leg.br\)](#) acesso em 08/04/2024 às 20:55.

SCHWAB, Klaus Schwab - **A Quarta Revolução Industrial**. São Paulo: Edipro. 2016. ISBN: 9788572839785, p.186.

Silveira, Alexandre Di Miceli. **Governança corporativa**: o essencial para líderes - 1. ed. - Rio de Janeiro: Elsevier, 2014. página 02.

UNIÃO EUROPEIA. Diretrizes éticas para uma IA de confiança (Ethics Guidelines For Trustworthy AI). Disponível em: [https://ec.europa.eu/digital-single-market/en/news/ethics-guidelines-trustworthy-ai]. Acesso em:18/06/2022

X2inteligênciadigital.Históriadainteligênciaartificial.X2inteligênciadigital,2022. Disponível em:https://x2inteligencia.digital/2020/02/20/historia-da-inteligenciaartificial-2/. Acesso em 15/04/2022.