



AI ACT EUROPEU E PL 2.883/23 A REGULAMENTAÇÃO DA INTELIGÊNCIA ARTIFICIAL ATRAVÉS DE CLASSIFICAÇÕES DE RISCOS

Ricardo Francisco de Lima Filho

Orientador: Prof. Dr. Celso Naoto Kashiura Jr.

Resumo: O presente trabalho pretende uma comparação entre o Artificial Intelligence Act, cujo texto foi proposto pela Comissão Europeia, órgão competente da União Europeia, em 2021 e o Projeto de Lei nº 2.883/23, proposto em maio de 2023 pelo Senador Rodrigo Pacheco (PSD/MG), fortemente inspirado no projeto europeu porém inserido num ordenamento jurídico completamente distinto e numa realidade socioeconômica também diferente da presente na Europa. O projeto europeu traz uma abordagem de regulamentações diferenciadas para cada uso de sistemas de inteligência artificial baseadas no grau de risco oferecido por cada hipótese, dividindo toda a gama de possibilidades em quatro categorias de risco, atribuindo restrições e normatizações distintas para cada um deles. O projeto brasileiro segue nessa toada, apresentando três graus de risco com normatizações distintas, com clara influência do projeto europeu mas adaptando as disposições à realidade brasileira. Foram comparadas as disposições de cada um dos projetos, bem como consideradas e apresentadas análises críticas de cada um dos projetos legislativos, com a finalidade de esclarecer o sistema que prevê cada um deles e destacar pontos onde o projeto europeu inspirou o brasileiro, mas principalmente onde o último inovou sobre seu antecessor.

Palavras-chave: Inteligência artificial. Projeto de Lei 2.883/23. Artificial Intelligence Act. Regulamentação. Inovação.

Abstract: This paper intends to compare the Artificial Intelligence Act, whose text was proposed by the European Commission, part of the European Union in 2021 and Bill n. 2.883/23, proposed by Senator Rodrigo Pacheco (PSD/MG) in 2023, deeply inspired by the European project but amidst a completely distinct legal system and socioeconomic reality. The European project brings a different approach to regulation for each use of artificial intelligence systems based on the levels of risk each hypothesis offers, dividing the whole sum of possibilities into four risk categories, attributing different restrictions and rules to each one. The Brazilian project follows in its stead, presenting three risk levels with differing restrictions and rules, showing clear influence from the European project but adapting the provisions to the Brazilian reality. The provisions of each project were compared to each other, and critical analyses of each one of them were considered and presented, with the intent of clarifying what regulatory system each project designs and highlighting points where the European project inspired the Brazilian one, but most of all, where the latter innovated on its antecessor.

Keywords: Artificial intelligence. Bill n° 2.883/23. Artificial Intelligence Act. Regulation. Innovation.

Introdução

Desde que a ideia de sistemas capazes de emular a inteligência do ser humano passou a habitar a imaginação humana — décadas antes de tal sistema ser sequer possível, já se imaginou que estes sistemas precisariam inevitavelmente ser regulamentados e limitados de alguma forma, sob pena de não só serem abusados por agentes de má-fé, como de agirem por conta própria, evoluindo além do controle humano.

Isaac Asimov, em 1950, publicou “Eu, Robô”, obra que revolucionou a ficção científica — que até então apresentava máquinas que se voltavam contra seus criadores, mas sem profundidade — ao apresentar a possibilidade de que a mente artificial poderia, um dia, atingir ou ultrapassar a complexidade das mentes naturais.

A obra, narrada por uma “*robopsicóloga*” que conta eventos marcantes de sua carreira, tornou-se um marco por apresentar a primeira proposta, mesmo que na ficção, de como regulamentar a inteligência artificial, através das Três Leis da Robótica:

- 1) um robô não pode ferir um humano ou permitir que um humano sofra algum mal;
- 2) os robôs devem obedecer às ordens dos humanos, exceto nos casos em que tais ordens entrem em conflito com a primeira lei;
- 3) um robô deve proteger sua própria existência, desde que não entre em conflito com as leis anteriores. (ASIMOV, 1950)

Em 1968, estreou o filme 2001: Uma Odisseia no Espaço, onde astronautas numa missão a caminho de Júpiter são traídos por HAL 9000, o computador que comanda os sistemas da aeronave *Discovery*, que descreve a si mesmo como “infalível e incapaz de erro”. Os astronautas, preocupados com as instruções que o computador lhes dá, decidem por desativá-lo, porém HAL 9000, sob a alegação de que ser desativado poria em risco a missão, mata um dos tripulantes e tenta matar o segundo, que consegue, enfim, desconectar os módulos de memória do computador.

O tema tem evoluído numa velocidade exponencial nos últimos anos, deixando de ser tema de ficção científica ou assunto de tratativas apenas por entusiastas e especialistas e passando a cada vez mais permear o cotidiano, onde o uso de sistemas de inteligência artificial se espalha por cada vez mais áreas da vida humana.

A partir disto, o presente artigo busca analisar a primeira proposta legislativa relevante neste sentido — o Artificial Intelligence Act da União Europeia — e compará-lo com o projeto recentemente apresentado no Senado Federal após longas deliberações e análises, que, apesar de fortemente inspirado no seu antecessor europeu, traz inovações e modificações bastante relevantes ao ordenamento jurídico pátrio.

Iniciar-se-á pela apresentação da inteligência artificial enquanto tecnologia de propósito geral, demonstrando a urgente necessidade de se regulamentar tais sistemas. Em seguida, será apresentado e analisado o Artificial Intelligence Act, bem como as perspectivas de

críticos e apoiadores do projeto. Ato contínuo, apresentar-se-á o Projeto de Lei nº 2.338/23, já delineando pontos de comparação entre os dispositivos, concluindo-se com uma análise de algumas das distinções e similaridades entre os projetos.

Por fim, o artigo propõe-se a, dada sua grande relevância, analisar brevemente uma tangente sobre as especificidades de legislações e regulamentações já existentes no Brasil e na União Europeia sobre o uso de sistemas de inteligência artificial no campo do direito e no Poder Judiciário.

I. Inteligência Artificial como Tecnologia de Propósito Geral

Antes de comparar os projetos de regulamentação de inteligência artificial e as normas propostas, é relevante discutir brevemente o conceito de inteligência artificial e seu impacto como tecnologia de propósito geral.

Embora a inteligência artificial, por si só, já seja um conceito bastante complexo e cuja definição exata é ainda relativamente controversa, ambos o AI Act e o Projeto de Lei 2.883/23 trazem definições aplicáveis no âmbito de suas respectivas abrangências.

O Projeto de Lei traz, em seu artigo 4º, inciso I, sua definição de ‘sistema de inteligência artificial’.

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

I – sistema de inteligência artificial: sistema computacional, com graus diferentes de autonomia, desenhado para inferir como atingir um dado conjunto de objetivos, utilizando abordagens baseadas em aprendizagem de máquina e/ou lógica e representação do conhecimento, por meio de dados de entrada provenientes de máquinas ou humanos, com o objetivo de produzir previsões, recomendações ou decisões que possam influenciar o ambiente virtual ou real; (BRASIL, 2023)

O conceito apresentado pelo projeto brasileiro traz um conceito bastante próximo ao adotado pelo projeto europeu, seu antecessor, limitando-se a esclarecer e especificar as técnicas aplicáveis e a influência no “ambiente virtual ou real”.

Artigo 3.º

Definições

Para efeitos do presente regulamento, entende-se por:

1) «Sistema de inteligência artificial» (sistema de IA), um programa informático desenvolvido com uma ou várias das técnicas e abordagens enumeradas no anexo I, capaz de, tendo em vista um determinado conjunto de objetivos definidos por seres humanos, criar resultados, tais como conteúdos, previsões, recomendações ou decisões, que influenciam os ambientes com os quais interage; (COMISSÃO Europeia, 2021)

A primeira referência ao conceito de inteligência artificial próxima ao que se tem hoje é atribuída a Alan Turing (KAUFMAN, 2019), matemático e cientista da computação inglês, que, em 1950 — mesmo ano em que “Eu, Robô” foi publicado, discorreu, com base nos conhecimentos de mais de sete décadas atrás, sobre o que seria possível com o passar do tempo e o avanço tecnológico.

Essa propriedade especial de computadores digitais, que podem imitar qualquer máquina de estados finita, é descrita dizendo que são máquinas universais. A

existência de máquinas com esta propriedade tem a consequência importante de que, considerações sobre velocidade a parte, é desnecessário criar várias máquinas novas para realizar vários processos computacionais. Todos eles podem ser feitos com um computador digital, programado de acordo com cada caso. Poderá ser visto que, como consequência disto, todos os computadores digitais serão de certa forma equivalentes. (TURING, 1950, p. 433)

Podemos esperar que máquinas eventualmente possam competir com homens em todos os campos puramente intelectuais. Mas quais são os melhores para começarmos? Até esta é uma decisão difícil. Muitos acreditam que uma atividade muito abstrata, como jogar xadrez, seria a melhor opção. Também há quem sustente que o melhor é fornecer à máquina os melhores sensores que o dinheiro pode comprar, e então ensiná-la a entender e falar inglês. Esse processo poderia seguir o ensino normal de uma criança. Coisas seriam apontadas e nomeadas, etc. Mais uma vez, não sei qual seria a resposta correta, mas acredito que ambas abordagens devam ser tentadas. Podemos ver apenas uma curta distância à nossa frente, mas já vemos muito a que se fazer. (TURING, 1950, p. 460)

Turing, em seu artigo, além de prever com precisão que o xadrez e a compreensão e produção de linguagem, especificamente, seriam pontos revolucionários para o campo da inteligência artificial, já ali falava sobre a criação de computadores como “máquinas universais”, capazes de serem programadas para atender à diversas finalidades.

Essa previsão se concretizou com o advento da computação, que, tida como ‘tecnologia de propósito geral’, revolucionou absolutamente o modo de vida humano.

Dora Kaufman, numa palestra realizada em 2023 sobre as implicações éticas e sociais da inteligência artificial, resumiu o conceito de tecnologia de propósito geral, considerando a IA como parte deste limitadíssimo rol de inovações.

A inteligência artificial é distinta, ela é uma tecnologia digital, mas é distinta das outras tecnologias digitais, porque ela é o que é considerado como tecnologia de propósito geral. Então, só para dar uma ideia, as últimas tecnologias consideradas de propósito geral foram o carvão, que iniciou a Revolução Industrial, a eletricidade e a computação.

O que caracteriza, pela definição, uma tecnologia de propósito geral?

Ela reconfigura, ela muda a lógica de funcionamento da economia e da sociedade. (ESCOBAR, 2023)

A pesquisadora expande a discussão sobre o impacto da IA em suas obras:

Na segunda década do século XXI, a convergência de diversas tecnologias tem promovido resultados superiores a quaisquer previsões precedentes (ainda que aquém da ficção científica). As máquinas e sistemas inteligentes estão executando tarefas que até recentemente eram prerrogativas dos humanos, e em alguns casos com resultados mais rápidos e mais assertivos. Mas é apenas uma década de “revolução”, e as máquinas ainda estão restritas a prever cenários (capacidade preditiva) com base em grandes conjuntos de dados e a executar tarefas específicas, sob a supervisão direta dos especialistas em ciência da computação. Esse relativamente pequeno avanço da IA, contudo, trouxe benefícios inéditos para a humanidade.

A inteligência artificial está transformando nossa relação com a tecnologia, e é a base da revolução digital em curso a partir da confluência de tecnologias do mundo digital

(internet das coisas/IoT, *blockchain*, plataformas digitais), do mundo físico (veículos autônomos, impressão 3D, robótica avançada, novos materiais) e do mundo biológico (manipulação genética). (KAUFMAN, 2019)

Este impacto grande, porém relativamente imprevisível, tido como capaz de alterar substancialmente a estrutura da sociedade como a conhecemos, é o motor por trás dos projetos de regulamentação ao redor do mundo, dentre os quais se destaca o AI Act (*Artificial Intelligence Act*), projeto legislativo publicado em 2021 e atualmente em debate na Comissão Europeia e com previsão de aprovação até o fim de 2023 (POUGET, 2023), não só por ter sido o primeiro modelo de regulamentação estatal da área a ser proposto e discutido e por ser bastante ambicioso (MENENGOLA, *et al*, 2023), mas também pela grande relevância do mercado da União Europeia, um dos maiores e mais importantes da economia mundial (RUSCHEMEIER, 2023).

II. O AI Act e o Sistema Regulatório da União Europeia

O AI Act europeu, apesar de possuir natureza de *regulation* (MENENGOLA, *et al*, 2023) dentre os tipos de legislação que a União Europeia é capaz de produzir, e, assim, ter caráter obrigatório e de aplicação direta nos Estados membros, utiliza o que é chamado de *requisitos essenciais* para estabelecer *benchmarks* obrigatórias para que produtos, serviços e outros usos de inteligência artificial acessem o mercado da União Europeia.

Numa síntese, tais requisitos determinam fins, mas não estabelecem obrigatoriedade de meios específicos para que os fins sejam atingidos, delineando características mínimas que poderão ser atingidas de várias formas distintas. Os *requisitos essenciais* são acompanhados das chamadas *normas harmonizadas*, que funcionam como recomendações ou sugestões de como atingir os objetivos delineados nas normas, normalmente, da forma mais direta e simplificada possível (UNIÃO EUROPEIA, 2023).

A observação das *normas harmonizadas*, embora não seja o único meio de atingir os requisitos essenciais, costuma ser o meio mais adequado (EBERS, 2021), uma vez que sua utilização gera presunção relativa de conformidade à lei e evita que se precise demonstrar e defender a eficácia das medidas alternativas escolhidas nas avaliações de conformidade às normas regulatórias, que poderão ser realizadas pelo próprio ente privado que pretende lançar um produto ou serviço ao mercado, como por comitês ou até diretamente por autoridades públicas, a depender do grau de risco do produto ou serviço analisado, sem prejuízo de fiscalização posterior (UNIÃO EUROPEIA, 2023).

Tanto os *requisitos essenciais* quanto as *normas harmonizadas* que os acompanham são delimitados propositalmente na maneira mais ampla e “não-técnica” possível, para que se abranja o maior número possível de diferentes casos e para que a norma não se torne obsoleta assim que novas formas e modelos de inteligência artificial sejam desenvolvidos (MCFADDEN, *et al*, 2021).

A partir deste sistema, o AI Act desenha um sistema de “níveis de riscos”, categorizando eventuais usos e propósitos para a inteligência artificial, baseando-se no nível de risco oferecido, dividindo os usos em categorias (FINOCCHIARO, 2023).

II.1. Risco Inaceitável

A primeira categoria, descrita no Título II, é a de risco inaceitável, delimitando hipóteses de uso da inteligência artificial que se consideram contrárias aos valores da União Europeia e violadores de direitos fundamentais, caracterizadas primariamente pela análise do fim pretendido (COMISSÃO EUROPEIA, 2021).

Tais usos são, portanto, proibidos e nenhum tipo de adequação aos *requisitos essenciais* torna-os viáveis para o acesso ao mercado europeu.

O primeiro destes fins se caracteriza pelo seguinte:

- a) A colocação no mercado, a colocação em serviço ou a utilização de um sistema de IA que empregue técnicas subliminares que contornem a consciência de uma pessoa para distorcer substancialmente o seu comportamento de uma forma que cause ou seja suscetível de causar danos físicos ou psicológicos a essa ou a outra pessoa;
- b) A colocação no mercado, a colocação em serviço ou a utilização de um sistema de IA que explore quaisquer vulnerabilidades de um grupo específico de pessoas associadas à sua idade ou deficiência física ou mental, a fim de distorcer substancialmente o comportamento de uma pessoa pertencente a esse grupo de uma forma que cause ou seja suscetível de causar danos físicos ou psicológicos a essa ou a outra pessoa;

A proibição constante do Artigo 5º, número 1, alínea ‘a’ não é absoluta no sentido de proibir quaisquer “técnicas subliminares” aplicadas ao usuário que pretendam “distorcer substancialmente seu comportamento”, mas proíbe especificamente aquelas que causem ou possam causar “danos físicos ou psicológicos a essa ou outra pessoa”.

Uma crítica levantada contra a redação do dispositivo é a adoção do termo “técnicas subliminares”, que não tem uma definição específica no texto nem em outras legislações da União Europeia, tornando a aplicabilidade da proibição para eventuais usos incerta. A depender da definição assumida, por exemplo, a utilização de inteligência artificial em algoritmos de pesquisa poderia ser proibida, uma vez que os resultados apresentados ao usuário podem, em teoria, alterar seu comportamento (BERMÚDEZ, *et al*, 2023).

Já a proibição da alínea ‘b’, apesar de similar à da alínea anterior, é mais específica, visando a proteção de pessoas consideradas vulneráveis.

O parágrafo 16 da exposição de motivos da proposta esclarece a proibição, explicitando seu caráter subjetivo e clarifica a não aplicabilidade da proibição para o que se chama de “efeitos legítimos”, ainda que apenas para pesquisa científica. O parágrafo, entretanto, não especifica o que seriam tais “efeitos legítimos”.

A intenção destes sistemas é distorcer substancialmente o comportamento de uma pessoa de uma forma que cause ou seja suscetível de causar danos a essa ou a outra pessoa. A intenção pode não ser detetada caso a distorção do comportamento humano resulte de fatores externos ao sistema de IA que escapam ao controlo do fornecedor ou do utilizador. A proibição não pode impedir a investigação desses sistemas de IA para efeitos legítimos, desde que essa investigação não implique uma utilização do sistema de IA em relações homem-máquina que exponha pessoas singulares a danos e seja efetuada de acordo com normas éticas reconhecidas para fins de investigação científica.

Já a alínea ‘c’ trata de sistemas de ranqueamento ou classificação social:

- c) A colocação no mercado, a colocação em serviço ou a utilização de sistemas de IA por autoridades públicas ou em seu nome para efeitos de avaliação ou classificação da credibilidade de pessoas singulares durante um certo período com base no seu comportamento social ou em características de personalidade ou pessoais, conhecidas ou previsíveis, em que a classificação social conduz a uma das seguintes situações ou a ambas:
 - i) tratamento prejudicial ou desfavorável de certas pessoas singulares ou grupos inteiros das mesmas em contextos sociais não relacionados com os contextos nos quais os dados foram originalmente gerados ou recolhidos,
 - ii) tratamento prejudicial ou desfavorável de certas pessoas singulares ou grupos inteiros das mesmas que é injustificado e desproporcionado face ao seu comportamento social ou à gravidade do mesmo;

De maneira semelhante à proibição das alíneas ‘a’ e ‘b’, esta também se baseia especificamente no resultado obtido pelo uso, ao invés de proibir o ranqueamento de indivíduos enquanto prática. Ressalte-se que a alínea ‘c’ não proíbe tais usos por agentes privados, apenas “por autoridades públicas ou em seu nome” (SMUHA, *et al*, 2021).

Ainda, relevante apontar que a proibição também não atinge classificações sociais cujos tratamentos sejam considerados justificados e proporcionais face ao critério que se utiliza para a classificação, mesmo que gerem tratamentos prejudiciais ou desfavoráveis para alguns indivíduos ou grupos. A norma, por sua vez, não traz critérios para que se determine o que tornaria um tratamento prejudicial justificado ou não.

Por fim, a alínea ‘d’, a mais específica e com uma lista de exceções de todas as proibições, trata de sistemas que utilizam inteligência artificial para identificar indivíduos em espaços públicos:

- d) A utilização de sistemas de identificação biométrica à distância em «tempo real» em espaços acessíveis ao público para efeitos de manutenção da ordem pública, salvo se essa utilização for estritamente necessária para alcançar um dos seguintes objetivos:
 - i) a investigação seletiva de potenciais vítimas específicas de crimes, nomeadamente crianças desaparecidas,
 - ii) a prevenção de uma ameaça específica, substancial e iminente à vida ou à segurança física de pessoas singulares ou de um ataque terrorista,
 - iii) a deteção, localização, identificação ou instauração de ação penal relativamente a um infrator ou suspeito de uma infração penal referida no artigo 2.º, n.º 2, da Decisão-Quadro 2002/584/JAI do Conselho e punível no Estado-Membro em causa com pena ou medida de segurança privativas de liberdade de duração máxima não inferior a três anos e tal como definidas pela legislação desse Estado-Membro.

Embora a norma traga um rol taxativo de casos específicos em que se pode usar inteligência artificial para identificação biométrica, ela traz elementos que devem ser considerados nesses usos, bem como a necessidade de obtenção de autorização para tal:

- (20) A fim de assegurar que esses sistemas sejam utilizados de uma forma responsável e proporcionada, também importa estabelecer que, em cada uma dessas três situações enunciadas exaustivamente e definidas de modo restrito, é necessário ter em conta determinados elementos, em especial no que se refere à natureza da situação que dá origem ao pedido e às consequências da utilização para os direitos e as liberdades de

todas as pessoas em causa e ainda às salvaguardas e condições previstas para a utilização. Além disso, a utilização de sistemas de identificação biométrica à distância «em tempo real» em espaços acessíveis ao público para efeitos de manutenção da ordem pública deve estar sujeita a limites espaciais e temporais adequados, tendo em conta, especialmente, os dados ou indícios relativos às ameaças, às vítimas ou ao infrator. A base de dados de pessoas utilizada como referência deve ser adequada a cada utilização em cada uma das três situações acima indicadas.

(21) Cada utilização de um sistema de identificação biométrica à distância «em tempo real» em espaços acessíveis ao público para efeitos de manutenção da ordem pública deve estar sujeita a uma autorização expressa e específica de uma autoridade judiciária ou de uma autoridade administrativa independente de um Estado-Membro. Em princípio, essa autorização deve ser obtida antes da sua utilização, salvo em situações de urgência devidamente justificadas, ou seja, quando a necessidade de utilizar os sistemas em causa torna efetiva e objetivamente impossível obter uma autorização antes de iniciar essa utilização. Nessas situações de urgência, a utilização deve limitar-se ao mínimo absolutamente necessário e estar sujeita a salvaguardas e condições adequadas, conforme determinado pelo direito nacional e especificado no contexto de cada caso de utilização urgente pela própria autoridade policial. Ademais, nessas situações, a autoridade policial deve procurar obter quanto antes uma autorização, apresentando as razões para não ter efetuado o pedido mais cedo.

O projeto, entretanto, enfrenta críticas pela amplitude das exceções previstas na proibição, uma vez que, por exemplo, não proíbe usos por agentes públicos fora do escopo da “manutenção da ordem pública”. Ainda, o uso para a busca de infratores ou suspeitos é amplo e coloca o poder discricionário para tal nas mãos das autoridades policiais e agências de inteligência locais. Por fim, considera-se que o mero fato de que a infraestrutura para tais usos existe e pode — ou não — estar em uso a qualquer momento pode prejudicar o exercício de direitos fundamentais, como os de liberdade de expressão, de manifestação e de associação, especialmente em se tratando de minorias (SMUHA, *et al*, 2021).

Críticos da proposta apontam, ainda, a exceção posta a todo uso tido como “desenvolvidos ou usados exclusivamente para fins militares”, conforme previsão expressa do Art. 2º, 3 do AI Act. A exceção é criticada por ser vaga ao não trazer qualquer condição ou previsão de possíveis usos militares tidos como aceitáveis, mas também por não trazer disposições sobre sistemas que forem desenvolvidos para fins militares mas venham a ser utilizadas em contextos civis ou vice-versa (RUSCHEMEIER, 2023).

II.2. Risco Elevado

A segunda categoria de risco é a de Risco Elevado, delineada no Título III da proposta, na qual estão previstos os usos que “só podem ser colocados no mercado da União [Europeia] ou colocados em serviço se cumprirem determinados requisitos obrigatórios” (COMISSÃO Europeia, 2021).

Há duas hipóteses para a classificação de um uso de inteligência artificial como sendo de risco elevado.

A primeira é o rol de casos constantes do Anexo III da proposta, que inclui hipóteses como a identificação e categorização biométrica de pessoas; recrutamento e seleção de

candidatos a vagas de emprego; migração e controle fronteiriço; segurança pública; o acesso a serviços públicos e privados, como prioridades para atendimentos médicos de urgência e emergência, elegibilidade para assistência pública ou pontuações usadas na concessão de crédito (COMISSÃO Europeia, 2021).

A segunda são os sistemas de inteligência artificial que são utilizados como componente de segurança de um produto, ou seja, ele próprio um produto abrangido pelas *normas harmonizadas* da União Europeia (COMISSÃO Europeia, 2021).

A norma conta com um vasto rol de exigências, incluindo sistemas de mitigação de danos, bancos de dados de alta qualidade que visem minimizar resultados discriminatórios, iteração continuada na análise de riscos, o registro de atividades para garantir a rastreabilidade dos resultados, documentação detalhada, garantia de informações claras e adequadas ao usuário, supervisão humana e um alto nível de robustez, segurança e precisão (COMISSÃO Europeia, 2021).

Dentre este rol de exigências, é relevante ressaltar o artigo 10, que trata dos bancos de dados de treino, validação e teste, bem como da governança e gestão destes dados, que devem ser pertinentes, representativos, completos e isentos de erros.

2. Os conjuntos de dados de treino, validação e teste devem estar sujeitos a práticas adequadas de governança e gestão de dados. Essas práticas dizem nomeadamente respeito:

- a) Às escolhas de conceção tomadas;
- b) À recolha de dados;
- c) Às operações de preparação e tratamento de dados necessárias, tais como anotação, rotulagem, limpeza, enriquecimento e agregação;
- d) À formulação dos pressupostos aplicáveis, nomeadamente no que diz respeito às informações que os dados devem medir e representar;
- e) À avaliação prévia da disponibilidade, quantidade e adequação dos conjuntos de dados que são necessários;
- f) Ao exame para detetar eventuais enviesamentos;
- g) À identificação de eventuais lacunas ou deficiências de dados e de possíveis soluções para as mesmas.

O projeto inova, nessa matéria, ao criar uma exceção ao disposto no GDPR (*General Data Protection Regulation* — ou Regulamento Geral sobre a Proteção de Dados, em português), regulamento europeu que versa sobre a privacidade e proteção de dados pessoais e proíbe o uso de informações pessoais tidas como sensíveis, como a etnia de um indivíduo, por exemplo. O AI Act, entretanto, no art. 10 (5), traz uma exceção à norma do GDPR ao permitir que tais informações sejam usadas em sistemas de inteligência artificial de alto risco — e apenas por seus provedores e desenvolvedores — para garantir que seus bancos de dados não contenham vieses discriminatórios. Apesar disso, o projeto não traz nenhuma exigência para mitigar riscos de vazamentos de tais dados sensíveis, nem sanções caso tais vazamentos venham a ocorrer (VEALE, BORGESIU, 2021).

Ainda, os provedores dos sistemas de inteligência artificial de risco alto são obrigados a construir um sistema chamado informalmente de “quatro olhos”, onde sistemas podem ser supervisionados por humanos, garantindo que riscos à saúde, segurança e outros direitos

fundamentais sejam minimizados. Deve haver um mínimo de dois supervisores humanos para que o sistema seja utilizado, que terão sua identidade registrada através de sistemas de identificação biométrica — garantindo assim, que existam ‘quatro olhos’ sobre o sistema antes de seu contato com o usuário final (VEALE, BORGESIU, 2021).

II.3. Risco Mínimo e Disposições Gerais

Todos os sistemas de inteligência artificial que não estiverem incluídos nas categorias de risco alto ou inaceitável — salvo os sujeitos a obrigações de transparência específicas, tratados no próximo item, serão considerados usos de IA de risco mínimo, não estando sujeitos a regulamentação específica e podendo ser utilizados e explorados livremente, desde que respeitem a legislação já existente.

Espera-se que a grande maioria dos usos atuais de inteligência estejam abarcados nesta categoria (COMISSÃO Europeia, 2023), que inclui usos como o gerenciamento de estoque em empresas, filtros de *spam* e a utilização em jogos, por exemplo.

Por tratar-se de uma categoria subsidiária às demais, que contam com um rol específico de usos previstos, não há previsão específica para os casos em que esta categoria seja aplicável, sendo necessário que os desenvolvedores que desejarem adaptar seus sistemas a esta categoria operem evitando o uso de componentes que elevem sua classificação de risco.

Apesar disso, o artigo 69 do projeto europeu determina que a Comissão Europeia e os Estados-Membros devem agir para incentivar que todos os sistemas de IA se adequem às especificações técnicas aplicáveis aos sistemas de risco elevado, porém sem que sejam aplicadas sanções aos que, não sendo obrigados a tal, não aderirem voluntariamente.

Devem ainda incentivar a adoção voluntária de requisitos como a sustentabilidade ambiental, a acessibilidade, a diversidade nas equipes de desenvolvimento, com a criação de indicadores de desempenho que permitam à Comissão Europeia e aos Estados-Membros medir o cumprimento de tais requisitos voluntários.

II.4. Obrigações de Transparência Aplicáveis a Determinados Sistemas

Além de todos os requisitos e obrigações a que estão obrigados os sistemas de inteligência artificial a depender da classificação de seu grau de risco, o AI Act, em seu artigo 52, prevê mais três hipóteses onde existirão obrigações adicionais relativas à transparência a serem cumpridas, sem qualquer prejuízo de outras dispostas na norma.

O primeiro item do artigo 52 trata de sistemas destinados a interagir com pessoas físicas, como *chatbots* de atendimento ao cliente, por exemplo, determinando que o sistema seja desenvolvido de uma forma que informe os usuários — explicitamente ou tornando óbvio pelo contexto — de que estão interagindo com uma inteligência artificial.

O segundo item trata de sistemas que objetivem o reconhecimento de emoções ou outras formas de categorização biométrica, devendo informar ao usuário não apenas de que estão interagindo com uma IA, mas informando também sobre seu funcionamento.

Já o terceiro item determina que sistemas que gerem ou manipulem imagens, áudios e/ou vídeos de forma que tais manipulações sejam verossímeis ou consideravelmente semelhantes a pessoas, objetos ou locais são obrigados a divulgar de maneira clara que o conteúdo foi gerado e/ou manipulado artificialmente, ressalvada hipótese onde tal uso viole a liberdade de expressão ou o direito à liberdade das artes e ciências, resguardados os direitos e liberdades de terceiros.

Nenhuma destas obrigações acima se aplica, entretanto, à sistemas que tenham como finalidade a prevenção, detecção, repressão ou investigação de crimes, desde que tais usos tenham autorização legal para funcionamento e não sejam disponibilizados ao público para recebimento de denúncias, no caso do item 1 (COMISSÃO Europeia, 2021).

Críticos, entretanto, questionam a efetividade de tais previsões, uma vez que, no caso de sistemas de reconhecimento de emoções, revisões de literatura recente concluem que não é possível identificar emoções de maneira confiável apenas através de análise de expressões faciais, bem como, em relação à IA generativa ou modificativa de imagens, áudio e/ou vídeo, ainda não existem métodos comprovados de identificação de tais atividades, tornando a fiscalização de seus usos extremamente difícil — senão impossível (VEALE, BORGESIU, 2021).

II.5. Sanções

As sanções previstas para descumprimento das determinações do AI Act são bastante similares às do GDPR (FLORIDI, *et al.*, 2022), prevendo a aplicação de *coimas* (ou multas, no português brasileiro).

Para entes privados, conforme previsão do artigo 71(3) do projeto, serão aplicadas multa de até trinta milhões de euros ou até seis por cento do faturamento bruto da empresa no exercício anterior, preferindo-se o maior valor, para violações que tratem do uso de sistemas de inteligência artificial tidos como de risco excessivo, e, portanto, proibidos; ou com a não conformidade com os requerimentos de governança de dados para usos tidos como de alto risco.

Outras violações da norma que não as descritas acima serão punidas com multa de até vinte milhões de euros ou até quatro por cento do faturamento bruto da empresa no ano anterior, preferindo-se o maior valor, conforme dispõe o artigo 71(4).

Já “O fornecimento de informações incorretas, incompletas ou enganadoras aos organismos notificados e às autoridades nacionais competentes em resposta a um pedido”, previsto no artigo 71(5), será punido com multa de até dez milhões de euros ou até dois por cento do faturamento bruto da empresa no ano anterior, preferindo-se o maior valor.

A fixação do valor da multa deve considerar “[a] natureza, a gravidade e a duração da infração e das suas consequências”, bem como eventual reincidência e a relevância do infrator no mercado.

O AI Act delega, através do artigo 71(7), a cada Estado-Membro da União Europeia a definição de normas sobre a aplicação de multas por eventuais infrações cometidas por autoridades e organismos de cada respectivo Estado-Membro.

Já infrações cometidas por instituições, órgãos e organismos da própria União Europeia

serão punidas com multas de até quinhentos mil euros no caso de uso de sistemas de inteligência artificial tidos como de risco excessivo e com a não conformidade com os requerimentos de governança de dados para usos tidos como de alto risco; e multas de duzentos e cinquenta mil euros para outras infrações.

Ainda, o AI Act prevê a existência de procedimento administrativo com garantia à ampla defesa, mesmo que tal previsão seja expressa apenas para instituições, órgãos e organismos da própria União Europeia (COMISSÃO Europeia, 2021).

Por fim, relevante ressaltar que todas estas sanções não trazem qualquer prejuízo ao direito de indivíduos que tenham sido diretamente impactados ou prejudicados por um sistema de inteligência artificial, que podem buscar jurisdição através do GDPR, para casos de violação de privacidade, por exemplo, ou por meio da reparação civil tradicional (FLORIDI, *et al.*, 2022).

III. O Projeto de Lei nº 2.338 de 2023

O Projeto de Lei nº 2338/23, protocolado pelo Senador Rodrigo Pacheco (PSD/MG) em 03 de maio de 2023 dispõe sobre o uso de sistemas de inteligência artificial no Brasil.

O texto, fortemente inspirado no AI Act Europeu e sua abordagem baseada em riscos e inovando em sua modelagem regulatória baseada em direitos (MENDONÇA JUNIOR, NUNES, 2023), bem como pensado em interação com a Lei Geral de Proteção de Dados (LGPD) brasileira (BRASIL, 2023), o projeto nasceu de um longo trabalho de pesquisa pela CJSUBIA (Comissão de Juristas Responsável por Subsidiar Elaboração de Substitutivo sobre Inteligência Artificial no Brasil), contando com participação popular em consultas e audiências públicas de diversos setores da sociedade (COALIZÃO, 2023).

A partir de uma abordagem regulatória baseada em riscos e em direitos (risks and rights-based approach), o texto do PL cria uma regulação assimétrica dos agentes regulados, com obrigações mais ou menos fortes de acordo com o nível de risco do sistema de IA, o que será determinado a partir de uma avaliação preliminar. Assim o PL estabelece direitos e medidas de governança básicos deflagrados por toda ferramenta de IA, mas também cria certos direitos e obrigações específicos para os casos potencialmente mais arriscados. Ao mesmo tempo, o projeto define que as medidas de governança dos sistemas de IA devem ser aplicadas ao longo de todo o seu ciclo de vida (desde a concepção até o seu encerramento/descontinuação).

O Projeto de Lei inova em relação ao projeto europeu ao prever expressamente, em seu art. 5º e seguintes, um rol de direitos específicos para pessoas afetadas por sistemas de inteligência artificial, enquanto seu antecessor limita-se a reafirmar a sujeição dos agentes aos direitos fundamentais já previstos na Carta de Direitos Fundamentais da União Europeia (COALIZÃO, 2023).

Art. 5º Pessoas afetadas por sistemas de inteligência artificial têm os seguintes direitos, a serem exercidos na forma e nas condições descritas neste Capítulo:

- I – direito à informação prévia quanto às suas interações com sistemas de inteligência artificial;
- II – direito à explicação sobre a decisão, recomendação ou previsão tomada por sistemas de inteligência artificial;

III – direito de contestar decisões ou previsões de sistemas de inteligência artificial que produzam efeitos jurídicos ou que impactem de maneira significativa os interesses do afetado;

IV – direito à determinação e à participação humana em decisões de sistemas de inteligência artificial, levando-se em conta o contexto e o estado da arte do desenvolvimento tecnológico;

V – direito à não-discriminação e à correção de vieses discriminatórios diretos, indiretos, ilegais ou abusivos; e

VI – direito à privacidade e à proteção de dados pessoais, nos termos da legislação pertinente.

Parágrafo único. Os agentes de inteligência artificial informarão, de forma clara e facilmente acessível, os procedimentos necessários para o exercício dos direitos descritos no caput.

Tais direitos são expandidos nos artigos seguintes, merecendo atenção especial o parágrafo único do art. 10 e o art. 11, que expandem a previsão do art. 5º, incisos III e IV.

Art. 10. Quando a decisão, previsão ou recomendação de sistema de inteligência artificial produzir efeitos jurídicos relevantes ou que impactem de maneira significativa os interesses da pessoa, inclusive por meio da geração de perfis e da realização de inferências, esta poderá solicitar a intervenção ou revisão humana.

Parágrafo único. A intervenção ou revisão humana não será exigida caso a sua implementação seja comprovadamente impossível, hipótese na qual o responsável pela operação do sistema de inteligência artificial implementará medidas alternativas eficazes, a fim de assegurar a reanálise da decisão contestada, levando em consideração os argumentos suscitados pela pessoa afetada, assim como a reparação de eventuais danos gerados.

Art. 11. Em cenários nos quais as decisões, previsões ou recomendações geradas por sistemas de inteligência artificial tenham um impacto irreversível ou de difícil reversão ou envolvam decisões que possam gerar riscos à vida ou à integridade física de indivíduos, haverá envolvimento humano significativo no processo decisório e determinação humana final.

Apesar da previsão do caput do artigo 10, que garante à pessoa o direito de solicitar intervenção ou revisão humana sempre que seus interesses forem significativamente impactados e/ou forem produzidos efeitos jurídicos relevantes, o parágrafo único do mesmo artigo e o art. 11 já preveem a possibilidade de tal intervenção ou modificação ser impossível, seja pela natureza da inteligência artificial utilizada — nos modelos de *blackbox* (BLOUIN, 2023), por exemplo — ou seja pela irreversibilidade dos efeitos causados, em ambos os casos deverão ser aplicadas medidas alternativas, como a reparação de eventuais danos e o envolvimento humano prévio à tomada de decisões (BRASIL, 2023).

Há de se mencionar ainda o artigo 18, que permite que a autoridade competente atualize a lista de sistemas de inteligência artificial e suas classificações de risco. A norma, entretanto, permite apenas a identificação de novas hipóteses, ou seja, permitindo-se a regulamentação de tipos não previstos em lei ou a majoração dos riscos de usos já previstos, mas não a redução do risco ou a exclusão de algum uso que conste expressamente dos rois de risco excessivo e alto.

III.1. Risco Excessivo

Além da previsão de direitos específicos, o Projeto de Lei propõe um sistema de riscos próximo ao do projeto europeu, iniciando por um rol similar de usos e finalidades de inteligência artificial vedados, por serem tidos como de risco excessivo:

Art. 14. São vedadas a implementação e o uso de sistemas de inteligência artificial:

I – que empreguem técnicas subliminares que tenham por objetivo ou por efeito induzir a pessoa natural a se comportar de forma prejudicial ou perigosa à sua saúde ou segurança ou contra os fundamentos desta Lei;

II – que explorem quaisquer vulnerabilidades de grupos específicos de pessoas naturais, tais como as associadas a sua idade ou deficiência física ou mental, de modo a induzi-las a se comportar de forma prejudicial a sua saúde ou segurança ou contra os fundamentos desta Lei;

III – pelo poder público, para avaliar, classificar ou ranquear as pessoas naturais, com base no seu comportamento social ou em atributos da sua personalidade, por meio de pontuação universal, para o acesso a bens e serviços e políticas públicas, de forma ilegítima ou desproporcional.

Art. 15. No âmbito de atividades de segurança pública, somente é permitido o uso de sistemas de identificação biométrica à distância, de forma contínua em espaços acessíveis ao público, quando houver previsão em lei federal específica e autorização judicial em conexão com a atividade de persecução penal individualizada, nos seguintes casos:

I – persecução de crimes passíveis de pena máxima de reclusão superior a dois anos;

II – busca de vítimas de crimes ou pessoas desaparecidas; ou

III – crime em flagrante.

Parágrafo único. A lei a que se refere o caput preverá medidas proporcionais e estritamente necessárias ao atendimento do interesse público, observados o devido processo legal e o controle judicial, bem como os princípios e direitos previstos nesta Lei, especialmente a garantia contra a discriminação e a necessidade de revisão da inferência algorítmica pelo agente público responsável, antes da tomada de qualquer ação em face da pessoa identificada.

Art. 16. Caberá à autoridade competente regulamentar os sistemas de inteligência artificial de risco excessivo.

O texto das proibições é bastante parecido com o previsto no AI Act, incluindo trechos copiados *ipsis litteris*, mas conta com algumas inovações relevantes.

A primeira é a adoção da expressão “ou contra os fundamentos desta Lei” nos incisos I e II do art. 14 para estender a proibição constantes destes a quaisquer fins que violem os fundamentos da norma, previstos nos incisos do artigo 2º:

Art. 2º O desenvolvimento, a implementação e o uso de sistemas de inteligência artificial no Brasil têm como fundamentos:

I – a centralidade da pessoa humana;

II – o respeito aos direitos humanos e aos valores democráticos;

III – o livre desenvolvimento da personalidade;

IV – a proteção ao meio ambiente e o desenvolvimento sustentável;

V – a igualdade, a não discriminação, a pluralidade e o respeito aos direitos trabalhistas;

VI – o desenvolvimento tecnológico e a inovação;

VII – a livre iniciativa, a livre concorrência e a defesa do consumidor;

- VIII – a privacidade, a proteção de dados e a autodeterminação informativa;
- IX – a promoção da pesquisa e do desenvolvimento com a finalidade de estimular a inovação nos setores produtivos e no poder público; e
- X – o acesso à informação e à educação, e a conscientização sobre os sistemas de inteligência artificial e suas aplicações.

A adoção desta expressão, combinada com o rol extenso de fundamentos do artigo 2º, além de expandir a proibição, torna as hipóteses de usos proibidos mais palpáveis e concretas (BRASIL, 2023).

A segunda distinção é a redação bem mais enxuta do art. 14, inciso III do PL 2.338/23, que, apesar de manter pontos do AI Act alvos de críticas, como a aplicabilidade da proibição apenas ao poder público e a inespecificidade dos termos “ilegítima” e “desproporcional”, o inciso é ainda menos abrangente, uma vez que vincula a proibição a utilização de “pontuação universal”, termo este que não encontra definição exata no dispositivo legal.

A terceira inovação em relação ao projeto europeu é a estrutura do art. 15, muito mais limitante que seu correspondente ao limitar a utilização de sistemas de identificação biométrica à distância apenas no cumprimento de três requisitos cumulativos, a previsão em lei federal, a autorização judicial e a previsão nos incisos do referido artigo. O AI Act, por sua vez, prevê a necessidade de autorização judicial, mas admite que tal autorização seja obtida durante ou após a utilização dos sistemas, sem que o uso não autorizado até então configure ilícito ou irregularidade.

Ainda no artigo 15, o parágrafo único prevê “a necessidade de revisão da inferência algorítmica pelo agente público responsável, antes da tomada de qualquer ação em face da pessoa identificada”, de forma a não dar soberania às decisões da inteligência artificial, garantindo que qualquer ação tomada a partir de sua utilização na identificação de indivíduos tenha passado por revisão humana pelo menos por um agente público.

Por fim, o artigo 16 atribui à autoridade competente regulamentar os sistemas de inteligência artificial tidos como de risco excessivo, posicionando as disposições desses artigos, assim, como um rol mínimo de normas, a partir das quais órgão ou entidade da Administração Pública Federal poderá expandir a regulamentação conforme se fizer necessário (BRASIL, 2023).

III.2. Alto Risco

Diferentemente do projeto europeu, o PL brasileiro traz um rol taxativo de finalidades para as quais o uso de sistemas de inteligência artificial é tido como de alto risco por seu potencial em afetar negativamente pessoas ou a sociedade e estará sujeito a um escrutínio maior e regulamentação mais rígida.

Art. 17. São considerados sistemas de inteligência artificial de alto risco aqueles utilizados para as seguintes finalidades:

- I – aplicação como dispositivos de segurança na gestão e no funcionamento de infraestruturas críticas, tais como controle de trânsito e redes de abastecimento de água e de eletricidade;

- II – educação e formação profissional, incluindo sistemas de determinação de acesso a instituições de ensino ou de formação profissional ou para avaliação e monitoramento de estudantes;
- III – recrutamento, triagem, filtragem, avaliação de candidatos, tomada de decisões sobre promoções ou cessações de relações contratuais de trabalho, repartição de tarefas e controle e avaliação do desempenho e do comportamento das pessoas afetadas por tais aplicações de inteligência artificial nas áreas de emprego, gestão de trabalhadores e acesso ao emprego por conta própria;
- IV – avaliação de critérios de acesso, elegibilidade, concessão, revisão, redução ou revogação de serviços privados e públicos que sejam considerados essenciais, incluindo sistemas utilizados para avaliar a elegibilidade de pessoas naturais quanto a prestações de serviços públicos de assistência e de segurança;
- V – avaliação da capacidade de endividamento das pessoas naturais ou estabelecimento de sua classificação de crédito;
- VI – envio ou estabelecimento de prioridades para serviços de resposta a emergências, incluindo bombeiros e assistência médica;
- VII – administração da justiça, incluindo sistemas que auxiliem autoridades judiciárias na investigação dos fatos e na aplicação da lei;
- VIII – veículos autônomos, quando seu uso puder gerar riscos à integridade física de pessoas;
- IX – aplicações na área da saúde, inclusive as destinadas a auxiliar diagnósticos e procedimentos médicos;
- X – sistemas biométricos de identificação;
- XI – investigação criminal e segurança pública, em especial para avaliações individuais de riscos pelas autoridades competentes, a fim de determinar o risco de uma pessoa cometer infrações ou de reincidir, ou o risco para potenciais vítimas de infrações penais ou para avaliar os traços de personalidade e as características ou o comportamento criminal passado de pessoas singulares ou grupos;
- XII – estudo analítico de crimes relativos a pessoas naturais, permitindo às autoridades policiais pesquisar grandes conjuntos de dados complexos, relacionados ou não relacionados, disponíveis em diferentes fontes de dados ou em diferentes formatos de dados, no intuito de identificar padrões desconhecidos ou descobrir relações escondidas nos dados;
- XIII – investigação por autoridades administrativas para avaliar a credibilidade dos elementos de prova no decurso da investigação ou repressão de infrações, para prever a ocorrência ou a recorrência de uma infração real ou potencial com base na definição de perfis de pessoas singulares; ou
- XIV – gestão da migração e controle de fronteiras.

Dentre todas estas hipóteses, merece destaque a previsão do inciso XI, que prevê a hipótese da utilização de vastos bancos de dados a fim de “avaliar os traços de personalidade e as características” de grupos determinados.

Em que pese a previsão se dar no contexto específico de investigações criminais e segurança pública e estar sujeita a uma gama de medidas de governança, não deixa de ser relevante a permissão a tais usos, mesmo que em caráter de alto risco, especialmente se consideradas as interações de tais sistemas de avaliação com os sistemas cujos usos estão previstos nos incisos III, VII, X, XIII, XIV.

Ainda, conforme proibição do artigo 14, III do PL 2.883/23, a utilização de qualquer tipo de classificação e/ou ranqueamento de indivíduos não pode ser considerada para os fins do inciso IV.

Como forma de mitigação destes riscos, o projeto obriga a adoção de medidas de governança específicas, em conjunto com as medidas já impostas a todos os sistemas de IA independentemente de sua classificação de risco.

Tais medidas são, em síntese, a documentação de todo o processo de desenvolvimento da tecnologia, incluindo decisões relevantes tomadas em sua construção e o fornecimento de informações ao operador e ao potencial afetado sobre o modelo, sua lógica e sua implementação e uso; o registro automático das operações do sistema de modo a demonstrar seu funcionamento e permitir o escrutínio e a mitigação de eventuais riscos; a realização de testes que avaliem a confiabilidade, robustez, a precisão e a acurácia, e a cobertura do sistema.

Além destas medidas, merece destaque a previsão do parágrafo único do art. 20, que determina a supervisão humana constante dos sistemas de inteligência artificial de alto risco para prevenir ou minimizar riscos para direitos e liberdades de pessoas que possam vir a ser afetadas pelo sistema (MENDONÇA JUNIOR, NUNES, 2023), exigindo ainda que os responsáveis pela supervisão humana em questão possam:

- I – compreender as capacidades e limitações do sistema de inteligência artificial e controlar devidamente o seu funcionamento, de modo que sinais de anomalias, disfuncionalidades e desempenho inesperado possam ser identificados e resolvidos o mais rapidamente possível;
- II – ter ciência da possível tendência para confiar automaticamente ou confiar excessivamente no resultado produzido pelo sistema de inteligência artificial;
- III – interpretar corretamente o resultado do sistema de inteligência artificial tendo em conta as características do sistema e as ferramentas e os métodos de interpretação disponíveis;
- IV – decidir, em qualquer situação específica, por não usar o sistema de inteligência artificial de alto risco ou ignorar, anular ou reverter seu resultado; e
- V – intervir no funcionamento do sistema de inteligência artificial de alto risco ou interromper seu funcionamento.

Por fim, entidades da Administração Pública Direta, além da avaliação de impacto algorítmico, estarão sujeitos à realização de consultas e audiências públicas, onde deverão ser divulgadas informações como origem dos dados utilizados, avaliações preliminares e garantias facilitadas de direito à explicação e revisão humana de decisões que gerem efeitos jurídicos consideráveis e/ou que afetem os interesses do indivíduo.

Se forem utilizados sistemas biométricos pela Administração Pública Direta, ainda deverá ser editado ato normativo que “estabeleça garantias para o exercício dos direitos da pessoa afetada e proteção contra a discriminação direta, indireta, ilegal ou abusiva, vedado o tratamento de dados de raça, cor ou etnia, salvo previsão expressa em lei” (MENDONÇA JUNIOR, NUNES, 2023).

III.3. Risco Mínimo e Medidas de Governança Amplas

Da mesma forma que no AI Act europeu, a categoria de Risco Mínimo é subsidiária, ou seja, caso um sistema de inteligência artificial não esteja abarcado pelo grau de risco excessivo ou alto risco, este sistema será automaticamente de risco mínimo, fazendo com que a grande maioria dos sistemas estejam incluídos nesta categoria.

Entretanto, independentemente do grau de risco em que um sistema de inteligência artificial for classificado no âmbito do Projeto de Lei 2.883/23, este estará sujeito a medidas de governança para que possa ser utilizado, com fim a garantir a segurança do sistema e o atendimento dos direitos de pessoas que possam ser por ele afetadas. Tais medidas de governança serão aplicáveis durante todo o ciclo de vida do sistema, incluindo durante sua concepção e desenvolvimento (BRASIL, 2023).

Tais medidas incluirão, nos termos do artigo 19 do projeto, medidas de gestão de dados para que se mitigue ou previna eventuais vieses discriminatórios; separação e organização dos dados destinados à treinamentos, testes e validações dos resultados do sistema; adoção de medidas de segurança da informação; adequação do tratamento de dados em conformidade com a legislação e a adoção de medidas de privacidade nos dados aplicáveis.

Deve ainda ser realizada a chamada Avaliação de Impacto Algorítmico, processo iterativo contínuo e público — respeitado o segredo industrial e comercial, onde serão considerados e registrados todos os riscos conhecidos e previsíveis do uso do sistema, considerando seus benefícios, eventuais consequências negativas e sua gravidade, a lógica de funcionamento do sistema, e serão previstas medidas de mitigação dos riscos e desenhados treinamentos e ações de conscientização para o uso do sistema (MACHADO, *et al.*).

Por fim, quando a sistemas de inteligência artificial que interajam diretamente com pessoas naturais, suas interfaces devem ser projetadas de forma que sejam adequadas e suficientemente claras e informativas — deixando claro ao usuário que este está interagindo com uma inteligência artificial.

III.4. Sanções e Responsabilidade Civil

O PL 2.883/23 adota um sistema de sanções bastante distinto do europeu, ao prever não apenas multas, mas um leque de possibilidades a serem aplicadas pela autoridade competente.

Art. 36. Os agentes de inteligência artificial, em razão das infrações cometidas às normas previstas nesta Lei, ficam sujeitos às seguintes sanções administrativas aplicáveis pela autoridade competente:

I – advertência;

II – multa simples, limitada, no total, a R\$ 50.000.000,00 (cinquenta milhões de reais) por infração, sendo, no caso de pessoa jurídica de direito privado, de até 2% (dois por cento) de seu faturamento, de seu grupo ou conglomerado no Brasil no seu último exercício, excluídos os tributos;

III – publicização da infração após devidamente apurada e confirmada a sua ocorrência;

IV – proibição ou restrição para participar de regime de *sandbox* regulatório previsto nesta Lei, por até cinco anos;

V – suspensão parcial ou total, temporária ou definitiva, do desenvolvimento, fornecimento ou operação do sistema de inteligência artificial; e

VI – proibição de tratamento de determinadas bases de dados.

Apesar da previsão dos incisos do *caput*, o parágrafo quarto do artigo 36 limita as penalidades a, no mínimo, as dos incisos II e V quando tratar-se de infração relativa ao desenvolvimento, fornecimento ou utilização de IA de risco excessivo.

O projeto brasileiro também expande o rol de critérios a serem considerados na fixação da penalidade e ao prever, expressamente, a necessidade de instauração de procedimento administrativo com o direito à ampla defesa antes da aplicação de qualquer sanção.

§ 1º As sanções serão aplicadas após procedimento administrativo que possibilite a oportunidade da ampla defesa, de forma gradativa, isolada ou cumulativa, de acordo com as peculiaridades do caso concreto e considerados os seguintes parâmetros e critérios:

I – a gravidade e a natureza das infrações e a eventual violação de direitos;

II – a boa-fé do infrator;

III – a vantagem auferida ou pretendida pelo infrator;

IV – a condição econômica do infrator;

V – a reincidência;

VI – o grau do dano;

VII – a cooperação do infrator;

VIII – a adoção reiterada e demonstrada de mecanismos e procedimentos internos capazes de minimizar riscos, inclusive a análise de impacto algorítmico e efetiva implementação de código de ética;

IX – a adoção de política de boas práticas e governança;

X – a pronta adoção de medidas corretivas;

XI – a proporcionalidade entre a gravidade da falta e a intensidade da sanção; e

XII – a cumulação com outras sanções administrativas eventualmente já aplicadas em definitivo para o mesmo ato ilícito.

Outra importante inovação é a previsão de medidas preventivas a serem adotadas antes ou durante o processo administrativo, com a possibilidade de aplicação de multas cominatórias, caso se verifiquem indícios de que o agente esteja causando ou possa vir a causar lesão irreparável ou de difícil reparação ou que este possa vir a tornar ineficaz o resultado do processo, ao apagar informações e destruir provas, por exemplo.

Por fim, o projeto avança também ao delimitar especificamente a responsabilidade civil por quaisquer danos causados pelo sistema de inteligência artificial que não sejam já abarcados pelo Código de Defesa do Consumidor por se tratarem de relações de consumo, aplicando a presunção de culpa — relativa ou absoluta, a depender do grau de risco — ao fornecedor e ao operador, permitindo, entretanto, a prova de culpa exclusiva da vítima e caso fortuito externo como forma de afastar sua responsabilidade.

Art. 27. O fornecedor ou operador de sistema de inteligência artificial que cause dano patrimonial, moral, individual ou coletivo é obrigado a repará-lo integralmente, independentemente do grau de autonomia do sistema.

§ 1º Quando se tratar de sistema de inteligência artificial de alto risco ou de risco excessivo, o fornecedor ou operador respondem objetivamente pelos danos causados, na medida de sua participação no dano.

§ 2º Quando não se tratar de sistema de inteligência artificial de alto risco, a culpa do agente causador do dano será presumida, aplicando-se a inversão do ônus da prova em favor da vítima.

Art. 28. Os agentes de inteligência artificial não serão responsabilizados quando:

I – comprovarem que não colocaram em circulação, empregaram ou tiraram proveito do sistema de inteligência artificial; ou

II – comprovarem que o dano é decorrente de fato exclusivo da vítima ou de terceiro, assim como de caso fortuito externo.

Art. 29. As hipóteses de responsabilização civil decorrentes de danos causados por sistemas de inteligência artificial no âmbito das relações de consumo permanecem sujeitas às regras previstas na Lei nº 8.078, de 11 de setembro de 1990 (Código de Defesa do Consumidor), sem prejuízo da aplicação das demais normas desta Lei.

IV. Distinções entre o Projeto de Lei 2.338/23 e o Artificial Intelligence Act

Apesar do Projeto de Lei brasileiro se inspirar fortemente no AI Act, estes possuem uma vasta gama de relevantes distinções que merecem destaque, a começar pela própria estrutura do ordenamento jurídico sobre a qual ambos os projetos foram desenhados e construídos.

O sistema proposto pelo projeto europeu baseia-se nas *normas harmonizadas*, requisitos amplos e relativamente genéricos a que os sistemas que desejem operar no mercado europeu devem se submeter, sem, entretanto, a exigência de que o façam na forma prevista na própria norma, desde que sejam capazes de comprovar que o objetivo foi alcançado e o requisito foi cumprido.

Já no projeto brasileiro, tentou-se adaptar o espírito da norma europeia ao ordenamento pátrio ao determinar um piso mínimo de medidas de governança delineadas na norma, que devem ser obrigatoriamente cumpridos, sem que fossem especificadas as ferramentas para tal, na maioria dos casos.

A estrutura da categoria de risco elevado ou alto é completamente diferente entre os projetos, embora seu conteúdo seja bastante similar, com o projeto europeu delimitando oito áreas onde a utilização de sistemas de inteligência artificial apresenta tal risco e limitando-se a trazer exemplos de casos onde o uso de tais sistemas será considerado pertencente a uma das áreas. Já o projeto brasileiro é mais conciso e específico, delineando finalidades específicas, cujo rol — embora possa ser expandido por autoridade competente — é taxativo.

O sistema de sanções também é bastante distinto. Enquanto o projeto europeu traz apenas sanções pecuniárias, o brasileiro expande o rol para prever diversos tipos de sanções a depender da gravidade da infração — partindo de simples advertências à proibição de utilizar determinadas bases de dados, incluindo a possibilidade de se aplicar multas cominatórias bem como a imposição de obrigação de reparar eventuais danos causados, sem prejuízo da penalidade pela infração.

O Projeto de Lei 2.883/23 preenche uma grande lacuna presente no texto europeu ao definir não só o conceito de discriminação, mas também o de discriminação indireta nos incisos VI e VII do artigo 4º, respectivamente.

A previsão específica e expressa da discriminação indireta é uma relevante medida no combate aos vieses algorítmicos, como no caso de sistemas de pontuação de crédito, que desproporcionalmente prejudicam populações periféricas e/ou racializadas, por exemplo (MACHADO, *et al.*, 2023).

VI – discriminação: qualquer distinção, exclusão, restrição ou preferência, em qualquer área da vida pública ou privada, cujo propósito ou efeito seja anular ou restringir o reconhecimento, gozo ou exercício, em condições de igualdade, de um ou mais direitos ou liberdades previstos no ordenamento jurídico, em razão de características pessoais como origem geográfica, raça, cor ou etnia, gênero, orientação sexual, classe socioeconômica, idade, deficiência, religião ou opiniões políticas;

VII – discriminação indireta: discriminação que ocorre quando normativa, prática ou critério aparentemente neutro tem a capacidade de acarretar desvantagem para pessoas pertencentes a grupo específico, ou as coloquem em desvantagem, a menos que essa normativa, prática ou critério tenha algum objetivo ou justificativa razoável e legítima à luz do direito à igualdade e dos demais direitos fundamentais;

Apesar dos avanços em relação ao projeto europeu, entretanto, o Projeto de Lei 2.883/23 acaba por omitir previsões relevantes, como a da avaliação do impacto dos sistemas de inteligência artificial em pequenas e médias empresas, que consta no projeto europeu mas não foi trazido para o brasileiro (SCHMIDT, 2023).

V. A Regulamentação da Inteligência Artificial no Direito

Apesar de avançar em diversos assuntos em relação a seu antecessor europeu, o Projeto de Lei brasileiro deixa de solucionar graves lacunas, principalmente ao tratar da previsão concernente ao uso de sistemas de inteligência artificial no Direito e na aplicação da lei, áreas extremamente sensíveis e com consequências na própria aplicação da norma aqui analisada, conforme reconhecido por instituições da União Europeia em mais de uma ocasião.

A primeira, editada pela CEPEJ, a Comissão Europeia Para a Eficácia da Justiça e intitulada “Carta Europeia de Ética sobre o Uso da Inteligência Artificial em Sistemas Judiciais e seu ambiente”, foi publicada em 2018 e tem, como principal objetivo, o estabelecimento de cinco princípios a serem observados por sistemas de IA que sejam utilizados no âmbito judicial:

Os cinco princípios da Carta Ética sobre o Uso da Inteligência Artificial nos Sistemas Judiciais e no respetivo ambiente

1 PRINCÍPIO DE RESPEITO AOS DIREITOS FUNDAMENTAIS: assegurar que a conceção e a aplicação de instrumentos e serviços de inteligência artificial sejam compatíveis com os direitos fundamentais.

2 PRINCÍPIO DE NÃO-DISCRIMINAÇÃO: prevenir especificamente o desenvolvimento ou a intensificação de qualquer discriminação entre indivíduos ou grupos de indivíduos.

3 PRINCÍPIO DE QUALIDADE E SEGURANÇA: em relação ao processamento de decisões e dados judiciais, utilizar fontes certificadas e dados intangíveis com modelos elaborados de forma multidisciplinar, em ambiente tecnológico seguro.

4 PRINCÍPIO DA TRANSPARÊNCIA, IMPARTIALIDADE E EQUIDADE: tornar os métodos de tratamento de dados acessíveis e compreensíveis, autorizar auditorias externas.

5 PRINCÍPIO "SOBRE O CONTROLO DO USUÁRIO": excluir uma abordagem prescritiva e garantir que os usuários sejam atores informados e controlem as escolhas feitas. (COMISSÃO, 2023)

O disposto na Carta, bastante compatível com os princípios constitucionais brasileiros (CASTRO, 2022), forma a base principiológica que deve ser observada sempre que se pretender a utilização de inteligência artificial nos sistemas judiciais.

Em uma segunda carta, esta editada pelo Conselho da União Europeia e titulada “A Carta dos Direitos Fundamentais no contexto da inteligência artificial e da transformação digital”, publicada em 2020 e tratando de outros temas além do uso da Inteligência Artificial:

Tecnologias digitais, incluindo IA, podem contribuir para melhorar o acesso à informação legal, possivelmente reduzindo a duração de procedimentos judiciais e melhorando o acesso à justiça em geral. Entretanto, tais desenvolvimentos também podem ter efeitos negativos, através do uso de algoritmos enviesados, por exemplo. Remédios legais efetivos devem ser garantidos para assegurar o direito a um julgamento justo, a presunção de inocência e o direito à defesa. Ainda, o acesso não-digital ao direito e à justiça continuará essencial. Nós continuamos comprometidos a defender e promover a garantia da lei na União e em seus Estados-Membros (CONSELHO, 2020).

O texto do AI Act atende a essa necessidade e prevê, no Anexo III, que traz o rol de casos tidos como de risco elevado, o item 8, (a), onde consta um único uso relativo à administração da justiça: “ (a) Sistemas de IA concebidos para auxiliar uma autoridade judiciária na investigação e na interpretação de *factos* e do direito e na aplicação da lei a um conjunto específico de *factos*” (COMISSÃO Europeia, 2021).

Esta redação foi duramente criticada por conter problemáticas importantes (SCHWEMER, *et al.*, 2021). A primeira trata da utilização do termo autoridade judiciária, que limita a classificação de risco elevado apenas para sistemas que auxiliem magistrados, fazendo com que sistemas que sejam concebidos para serem utilizados por advogados, promotores e outros operadores do direito — mesmo que magistrados venham a adotar sua utilização — sejam tidos como de risco mínimo, caso não estejam abarcadas por outras provisões.

Trata-se da mesma ambiguidade existente na exceção concedida pelo AI Act aos sistemas de uso militar. A previsão do item 8, (a) traz que “[s]istemas de IA concebidos para auxiliar uma autoridade judiciária” serão submetidos ao regime de risco elevado, ou seja, a determinação do grau de risco se dá na concepção do sistema e não em sua efetiva utilização e geração de potenciais efeitos.

Outra crítica se refere à redação do uso: “na investigação e na interpretação de *factos* e do direito e na aplicação da lei a um conjunto específico de *factos*”. O uso dos conectivos “e” entre as hipóteses gera insegurança jurídica, uma vez que não é possível precisar se, para que se qualifique para o risco elevado, o sistema precisa, cumulativamente, ser utilizado para investigação, interpretação e aplicação da lei, ou se apenas o enquadramento em uma das hipóteses é suficiente (SCHWEMER, *et al.*, 2021).

O parágrafo 40 da exposição de motivos do projeto, mesmo possuindo apenas caráter persuasivo na interpretação do texto legal, ajuda a elucidar a *mens legis* por trás da disposição ao afastar a aplicabilidade do item 8, (a) de atividades tidas como puramente auxiliares:

Contudo, essa classificação não deve ser alargada aos sistemas de IA concebidos para atividades administrativas puramente auxiliares que não afetam a administração efetiva da justiça em casos individuais, como a anonimização ou a pseudonimização de decisões judiciais, documentos ou dados, comunicações entre pessoal, tarefas administrativas ou afetação de recursos. (COMISSÃO Europeia, 2021)

O projeto brasileiro avança em relação ao texto europeu neste tema. O artigo 17, inciso VII do Projeto de Lei 2.883/23 prevê:

Art. 17. São considerados sistemas de inteligência artificial de alto risco aqueles utilizados para as seguintes finalidades:

(...)

VII – administração da justiça, incluindo sistemas que auxiliem autoridades judiciárias na investigação dos fatos e na aplicação da lei;

Apesar de ainda utilizar a expressão “autoridades judiciárias”, a nova redação modifica o momento onde se analisa a finalidade do sistema para sua efetiva utilização, bem como utiliza a expressão “incluindo”, demonstrando tratar-se de um exemplo de uso considerado abarcado em “administração da justiça”. Assim, o texto atende às principais críticas, mas sem desvirtuar a previsão da qual originou.

A previsão se une à Resolução 332/20 do Conselho Nacional de Justiça (CNJ), que dispõe sobre o uso de IA no Poder Judiciário, em específico sobre a ética, a transparência e a governança na produção e uso de tais sistemas. A resolução do CNJ, que expressamente menciona a Carta da CEPEJ como inspiração, reitera a base principiológica adotada pela União Europeia e a expande, adicionando ao rol de princípios a pesquisa e desenvolvimento, bem como a prestação de contas (BRASIL, 2020).

A resolução, entretanto, parece não ter sido capaz de, por si só, regulamentar o uso da inteligência artificial no Poder Judiciário.

Conforme levantamento realizado pelo próprio Conselho Nacional de Justiça em 2022, dos 111 projetos, espalhados por 53 tribunais, um número preocupante está em desacordo com o disposto na resolução. Importante destacar que o número de projetos com inteligência artificial no Judiciário Brasileiro saltou 171% de 2021 para 2022 (BRASIL, 2022).

Tratando de transparência, prevista no Capítulo IV da Resolução 332/20, 61 dos projetos não são ou não sabem dizer se são de código aberto, não sendo possível sua revisão externa. Tal situação preocupa, uma vez que do universo de 111 projetos, 99 declaram afetar mais de 1000 processos judiciais (BRASIL, 2022).

Ainda conforme o mesmo levantamento, em 25 dos 111 projetos a equipe técnica da instituição afirmou não ser capaz de explicar “o processo pelo qual entradas se tornam saídas”, ou seja, como as informações inseridas no sistema são tratadas e analisadas para gerar os resultados.

Mais grave, do total, apenas 61 dos projetos declaram ter passado por “monitoramento técnico e processos de garantia de qualidade”, 39 por “revisão legal e/ou administrativa” e 37 declaram não possuir documentação.

Por fim, em absoluto desacordo com a disposição do artigo 7º da Resolução 332/20, em especial o parágrafo 1º, apenas 38 projetos — pouco mais de um terço do total — declararam terem passado por “revisão de seus dados de treinamento para detectar vieses”.

CAPÍTULO III

DA NÃO DISCRIMINAÇÃO

Art. 7º As decisões judiciais apoiadas em ferramentas de Inteligência Artificial devem preservar a igualdade, a não discriminação, a pluralidade e a solidariedade, auxiliando no julgamento justo, com criação de condições que visem eliminar ou minimizar a opressão, a marginalização do ser humano e os erros de julgamento decorrentes de preconceitos.

§ 1º Antes de ser colocado em produção, o modelo de Inteligência Artificial deverá ser homologado de forma a identificar se preconceitos ou generalizações influenciaram seu desenvolvimento, acarretando tendências discriminatórias no seu funcionamento.

§ 2º Verificado viés discriminatório de qualquer natureza ou incompatibilidade do modelo de Inteligência Artificial com os princípios previstos nesta Resolução, deverão ser adotadas medidas corretivas.

§ 3º A impossibilidade de eliminação do viés discriminatório do modelo de Inteligência Artificial implicará na descontinuidade de sua utilização, com o consequente registro de seu projeto e as razões que levaram a tal decisão. (BRASIL, 2020)

Este levantamento demonstra a absoluta necessidade de regulamentação mais rígida e específica sobre o tema, especialmente num contexto tão sensível como é o Poder Judiciário.

Aprovado o Projeto de Lei 2.883/23, será reforçado o corpo regulatório, e, conseqüentemente, espera-se que as futuras iniciativas, bem como as já em andamento sejam aprimoradas de forma a se adequarem às novas normas.

Conclusões

Espera-se que ambos os projetos sejam aprovados brevemente, uma vez que frutos de longas discussões por parte das sociedades e das indústrias que afetam. Estes projetos representam uma oportunidade importante para o desenho de um sistema que considere as particularidades de cada sistema e impeça — ou pelo menos mitigue — prejuízos e eventuais abusos por agentes de má-fé.

Ambos os projetos, dentro de suas particularidades e apesar de ainda terem lacunas importantes, constroem estruturas que prezam pela garantia dos direitos fundamentais dos indivíduos que possam vir a ser afetados pelos sistemas de inteligência artificial, pela garantia da responsabilização de maus agentes e pela construção de boas práticas a serem adotadas preventivamente, até pelos sistemas não sujeitos a elas.

Entretanto, nenhum dos projetos, se aprovados, resolverá o problema. É essencial que os especialistas das mais diversas áreas e em tais sistemas, juristas, sociólogos e a sociedade como um todo entendam a importância e o impacto que o advento da inteligência artificial terá na sociedade que hoje conhecemos, mas, mais importante, que não permitam que o debate fique

confinado à tecnocracia da otimização, do avanço desenfreado, e da incessante mercantilização dos menores detalhes da vida humana. O debate deve seguir qualitativo, prezando, acima do atendimento aos interesses corporativos, pela proteção dos indivíduos, do coletivo, da privacidade e da sustentabilidade, postas em xeque com o avanço grandemente desregulado destas úteis, porém perigosas tecnologias.

Referências bibliográficas

- ACCESS Now Europe. **Access Now's submission to the European Commission's adoption consultation on the Artificial Intelligence Act**. Bruxelas, Bélgica: Access Now Europe, 2021.
- AS NORMAS na Europa. **União Europeia**, 2023. Disponível em https://europa.eu/youreurope/business/product-requirements/standards/standards-in-europe/index_pt.htm. Acesso em 26 ago. 2023.
- ASIMOV, Isaac. **Eu, robô**. São Paulo: Aleph, 2014.
- BERMÚDEZ, Juan Pablo; et al. What Is a Subliminal Technique? An Ethical Perspective on AI-Driven Influence. **2023 IEEE International Symposium On Ethics In Engineering, Science, And Technology (Ethics)**, [S.L.], p. 1-10, 18 maio 2023.
- BLOUIN, Lou. AI's mysterious 'black box' problem, explained. **University of Michigan-Dearborn**. 6 mar. 2023. Disponível em: <https://umdearborn.edu/news/ais-mysterious-black-box-problem-explained>. Acesso em 20 set. 2023.
- BRASIL. Autoridade Nacional de Proteção de Dados. **Análise preliminar do Projeto de Lei nº 2338/23, que dispõe sobre o uso da Inteligência Artificial**. Disponível em: <https://www.gov.br/anpd/pt-br/assuntos/noticias/analise-preliminar-do-pl-2338-2023-formatado-ascom.pdf>. Acesso em 01 out. 2023.
- BRASIL. Conselho Nacional de Justiça. **Resolução nº 332 de 21/08/2020**. Dispõe sobre a ética, a transparência e a governança na produção e no uso de Inteligência Artificial no Poder Judiciário e dá outras providências. Disponível em: <https://atos.cnj.jus.br/files/original191707202008255f4563b35f8e8.pdf>. Acesso em 02 out. 2023.
- BRASIL. Conselho Nacional de Justiça. **Resultados Pesquisa IA no Poder Judiciário - 2022**. Disponível em: <https://www.cnj.jus.br/sistemas/plataforma-sinapses/paineis-e-publicacoes/>. Acesso em 04 out. 2023.
- BRASIL. Senado Federal. **Projeto de Lei nº 2.338, de 3 de maio de 2023**. Dispõe sobre o uso da Inteligência Artificial. Brasília: Senado Federal, 2023. Disponível em: <https://www25.senado.leg.br/web/atividade/materias/-/materia/157233>. Acesso em: 01 set. 2023.
- CASTRO, Kátia Shimizu de. Princípios éticos Europeus no uso da inteligência artificial e a correlação com os Princípios Constitucionais Brasileiros. **Direito Internacional e Globalização Econômica (DIGE)**, São Paulo, v. 9, n. 9, p. 319-338, 04 ago. 2022.

COALIZÃO Direitos na Rede. **Carta de Apoio ao PL 2338/2023**, 2023. Disponível em: <https://direitosnarede.org.br/2023/06/14/carta-de-apoio-ao-pl-2338-2023/>. Acesso em 18 set. 2023.

COMISSÃO Europeia. **ANNEXES to the Proposal for a Regulation of the European Parliament and of the Council**. Bruxelas, Bélgica. 2021. Disponível em: https://eur-lex.europa.eu/resource.html?uri=cellar:e0649735-a372-11eb-9585-01aa75ed71a1.0001.02/DOC_2&format=PDF. Acesso em: 21 set. 2023.

COMISSÃO Europeia. **Charter of fundamental rights of the European Union**. Bruxelas, Bélgica, 2012. Disponível em: <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:12012P/TXT>. Acesso em: 20 set. 2023.

COMISSÃO Europeia. **Regulation of the European Parliament and of the Council laying down harmonized rules on artificial intelligence (artificial intelligence act) and amending certain union legislative acts**. Bruxelas, Bélgica. 2021. Disponível em: <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:52021PC0206>. Acesso em: 20 set. 2023.

COMISSÃO Europeia. **Regulatory framework proposal on artificial intelligence**. 20 jun. 2023. Disponível em: <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>. Acesso em 27 set. 2023.

COMISSÃO Europeia Para a Eficácia da Justiça (CEPEJ). **Carta Europeia de Ética sobre o Uso da Inteligência Artificial em Sistemas Judiciais e seu ambiente**. Estrasburgo, 2018. Disponível em: <https://rm.coe.int/carta-etica-traduzida-para-portugues-revista/168093b7e0>. Acesso em 09 out. 2023.

CONFORMITY Assessment. **Internal Market, Industry, Entrepreneurship and SMEs**, 2023. Disponível em: https://single-market-economy.ec.europa.eu/single-market/goods/building-blocks/conformity-assessment_en. Acesso em 25 ago. 2023.

CONSELHO da União Europeia. **A Carta dos Direitos Fundamentais no contexto da inteligência artificial e da transformação digital**. Bruxelas, Bélgica: Conclusões da Presidência, 2020.

EBERS, Martin. Standardizing AI - The Case of the European Commission's Proposal for an Artificial Intelligence Act. **SSRN Electronic Journal**, [S.L.], ago. 2021.

ESCOBAR, Herton. Inteligência artificial reconfigura a lógica de funcionamento da sociedade. **Jornal da USP**. São Paulo, 20 de jun. de 2023. Disponível em: <http://www.saocarlos.usp.br/inteligencia-artificial-reconfigura-a-logica-de-funcionamento-da-sociedade/>. Acesso em 18 de set. de 2023.

FINOCCHIARO, Giusella. The regulation of artificial intelligence. **Ai & Society**, [S.L.], v. 34, n. 4, p. 1-8, 3 abr. 2023. Disponível em: <https://link.springer.com/article/10.1007/s00146-023-01650-z>. Acesso em: 20 set. 2023.

FLORIDI, Luciano; et al. CapAI - A Procedure for Conducting Conformity Assessment of AI Systems in Line with the EU Artificial Intelligence Act. **SSRN Electronic Journal**, [S.L.], 23 mar. 2022.

- KAUFMAN, Dora. **A inteligência artificial irá suplantará a inteligência humana?**. 2ª ed., Barueri, Estação das Letras e Cores, 2019.
- KONIAKOU, Vasiliki. From the “rush to ethics” to the “race for governance” in Artificial Intelligence. **Information Systems Frontiers**, [S.L.], v. 25, n. 1, p. 71-102, 28 jun. 2022.
- MACHADO, Joana, *et al.* Sistemas de Inteligência Artificial e Avaliações de Impacto para Direitos Humanos. **Revista Culturas Jurídicas**. *Ahead of print*, 2023.
- MCFADDEN, Mark; JONES, Kate; TAYLOR, Emily; OSBORN, Georgia. **Harmonising Artificial Intelligence: The role of standards in the EU AI regulation**. Oxford, Inglaterra: University of Oxford, 2021.
- MENDONÇA JUNIOR, Claudio do Nascimento; NUNES, Dierle José Coelho. Desafios e Oportunidades para a Regulamentação da Inteligência Artificial: a necessidade de compreensão e mitigação dos riscos da IA. **Revista Contemporânea**, [S.L.], v. 3, n. 07, p. 7753-7785, 10 jul. 2023.
- MENENGOLA, Everton; GABARDO, Emerson; SANMIGUEL, Nancy Nelly González. The proposal of the european regulation on artificial intelligence. **Seqüência Estudos Jurídicos e Políticos**, Florianópolis, v. 43, n. 91, 1 fev. 2023.
- POUGET, Hadrien. Institutional Context. **The Artificial Intelligence Act Newsletter**. Disponível em <https://artificialintelligenceact.eu/context/>. Acesso em 28 de ago. de 2023.
- RUSCHEMEIER, Hannah. AI as a challenge for legal regulation - the scope of application of the artificial intelligence act proposal. **Era Forum**, [S.L.], v. 23, n. 3, p. 361-376, 9 jan. 2023.
- SCHMIDT, Sarah. Os desafios para regulamentar o uso da inteligência artificial. **Revista Pesquisa FAPESP**. Edição 331, set. 2023. Disponível em: <https://revistapesquisa.fapesp.br/os-desafios-para-regulamentar-o-uso-da-inteligencia-artificial/>. Acesso em 03 out. 2023.
- SCHWEMER, Sebastian Felix; TOMADA, Letizia; PASINI, Tommaso. Legal AI Systems in the EU’s proposed Artificial Intelligence Act. **Second International Workshop on AI and Intelligent Assistance for Legal Professionals in the Digital Workplace (LegalAIIA 2021)**. São Paulo, 21. jun. 2021.
- SHATKOVSKAYA, Tatiana V. et al. Artificial Intelligence Technology as a Complex Object of Intellectual Rights. **Technological Trends In The Ai Economy**, [S.L.], p. 159-168, 2023.
- SMUHA, N. A. et al. How the EU can achieve legally trustworthy AI: a response to the European Commission’s Proposal for an Artificial Intelligence Act. **SSRN Electronic Journal**, [S. L.], 5 ago. 2021.
- TURING, Alan M. Computing machinery and intelligence. **Mind**, vol. 59, no. 236, p. 433-460, 1950.
- VEALE, Michael; BORGESIU, Frederik Zuiderveen. Demystifying the Draft EU Artificial Intelligence Act — Analysing the good, the bad, and the unclear elements of the proposed approach. **Computer Law Review International**, [S.L.], v. 22, n. 4, p. 97-112, 1 ago. 2021.