

ços tecnológicos assim o permitirem, à conduta do *robot*. O caso da lesão provocada por *robot* autónomo configura um problema novo, que reclama, no mínimo, a adaptação das soluções jurídicas existentes, senão uma solução inteiramente nova, que permita ressarcir os danos decorrentes dessa lesão.

O EQUILÍBRIO TÊNUE ENTRE O PROGRESSO CIENTÍFICO E OS DIREITOS FUNDAMENTAIS: OS DANOS CAUSADOS POR IA À LUZ DA RESOLUÇÃO DO PARLAMENTO E DA PROPOSTA DE DIRETIVA

MARTA CANSADO²⁸⁶

SUMÁRIO: 1. Introdução. 2. A necessidade histórica do regime de responsabilidade objetiva. 3. As características particulares da Inteligência Artificial. 3.1. A falta de explicabilidade. 3.2. Enviesamentos e discriminação. 4. O regime de responsabilidade civil por danos causados por sistemas de IA. 4.1. Enquadramento. 4.2. A Proposta da União Europeia (Análise da Resolução do Parlamento Europeu de 2020 e Diretiva do Parlamento Europeu e do Conselho relativa à adaptação das regras de responsabilidade civil extracontratual à inteligência artificial). 4.2.1. A distinção com base no risco. 4.2.2. O regime de responsabilidade objetiva para os sistemas de IA de risco elevado. 4.2.3. O regime de responsabilidade culposa com presunção de culpa para os restantes sistemas de IA. 4.2.4. O regime de responsabilidade culposa e a presunção de nexo de causalidade da Proposta de Diretiva. 4.2.5. O dever de prestar informações. 5. Conclusão.

²⁸⁶ A autora é finalista da Licenciatura em Direito, Faculdade de Direito da Universidade Católica Portuguesa (Lisboa) e Investigadora Júnior na Revista “Vere Dictum Binário”, Cavaleiro & Associados, Sociedade de Advogados, R.L. Para além de ser Bolseira da Fundação Calouste Gulbenkian – Bolsa Novos Talentos –, foi estagiária de verão na Abreu Advogados e na CCA Law Firm. Foi oradora da Equipa Nacional Vencedora da Philip C. Jessup International Law Moot Court Competition, vencedora da Essay Competition “Implicações do Metaverso para o Direito: Como garantir uma transição sustentável para o mundo digital?”, org. European Law Student’s Association e Abreu Advogados, e vencedora dos Prémios Doutor CASTRO MENDES, e Doutor MANUEL GOMES DA SILVA.

RESUMO: O desenvolvimento da Inteligência Artificial tem vindo a revelar as vantagens, e também os riscos que estas tecnologias comportam. Para acautelar os danos provocados por sistemas de IA, discute-se atualmente qual o regime de responsabilidade civil que se afigura mais apropriado. Serão analisadas as propostas da UE nesta matéria, plasmadas na Resolução do Parlamento de 2020, e na Proposta de Diretiva de 2022 do Parlamento e do Conselho, relativas à responsabilidade civil aplicável aos danos causados por IA.

ABSTRACT: The development of Artificial Intelligence has revealed its advantages, as well as the risks it represents. In order to compensate for the injuries caused by AI systems, there is a discussion regarding the adequacy of the current civil liability regimes. The proposals of the EU will be discussed, with an analysis of the 2020 Resolution of the Parliament, and the 2022 Draft AI Liability Directive, which regard the civil liability regime applicable to injuries caused by AI systems.

PALAVRAS-CHAVE: Inteligência Artificial; Responsabilidade civil.

KEYWORDS: Artificial Intelligence; Civil Liability.

1. Introdução

O tópico da Inteligência Artificial (“IA”) consiste agora num dos temas mais discutidos, em várias áreas, particularmente na área do Direito. Enquanto algumas pessoas veem a Inteligência Artificial como o precursor de uma Quarta Revolução Industrial, outros poderão questionar se verdadeiramente, estaremos perante um avanço tecnológico tão significativo, que justifique o estabelecimento de um novo marco histórico.

A previsão de uma Quarta Revolução Industrial parece fundamentada, dado que a Inteligência Artificial tem revelado já a sua capacidade para substituir certos postos de trabalho, demonstrando, em muitos casos, uma maior eficiência e precisão na realização dos mesmos. Simultaneamente, as faculdades da Inteligência Artificial expandem-se para além do mero aperfeiçoamento do processo produtivo, pois a sua utilização é já equacionada em setores como a medicina, a administração pública, ou a justiça²⁸⁷.

²⁸⁷ PATRÃO, M. NEVES; ALMEIDA, A. BETÂMIO (2023). *Before and Beyond*

Por outro lado, as enormes potencialidades da IA são contrabalançadas com riscos significativos. Quer as características intrínsecas dos sistemas de IA, quer o modo como estes são aplicados, poderá gerar danos consideráveis, como inclusive já se tem vindo a demonstrar. Por esta razão, o Direito apercebeu-se da necessidade da sua intervenção, não só no sentido de prevenir a verificação destes danos, como para determinar a forma de estes serem compensados. A respeito da prevenção dos danos, foi já dado um passo importante, com a aprovação do Regulamento de Inteligência Artificial (*AI Act*), que estabelece obrigações para os produtores de sistemas de IA. No entanto, a questão relativa à compensação dos danos, em particular, quanto ao regime de responsabilidade civil concretamente

Artificial Intelligence: Opportunities and Challenges. In “Multidisciplinary Perspectives on Artificial Intelligence and the Law”, eds. HENRIQUE SOUSA ANTUNES, PEDRO FREITAS, ARLINDO OLIVEIRA, CLARA MARTINS PEREIRA, ELSA VAZ DE SEQUEIRA & LUÍS BARRETO XAVIER, Springer, Law, Governance and Technology Series, vol. 58, pp. 107-125.

te aplicável, encontra-se ainda em aberto.

Dentro deste tópico, importará destacar as intervenções legislativas da União Europeia, na forma de uma Resolução do Parlamento Europeu (Resolução do Parlamento Europeu, de 20 de outubro de 2020, que contém recomendações à Comissão sobre o regime de responsabilidade civil aplicável à inteligência artificial (2020/2014(INL)), bem como da Proposta de Diretiva Parlamento e do Conselho relativa à adaptação das regras de responsabilidade civil extracontratual à inteligência artificial (Diretiva Responsabilidade da IA).

Neste artigo, será primeiramente sustentada a necessidade de aplicação de um regime de responsabilidade objetiva, para certos sistemas de IA, fundamentando esta necessidade por comparação com outros regimes de responsabilidade objetiva já existentes. Num segundo plano, será avaliada a proposta de regime de responsabilidade civil, para danos causados por IA, resultante da Resolução e da Proposta de Diretiva, realizando uma análise crítica da mesma.

Em última análise, pretende-se contribuir para o desenvolvimen-

to de um regime de responsabilidade civil, para os danos causados por IA, que garanta simultaneamente o incentivo necessário para o desenvolvimento da Inteligência Artificial, com todas as suas vantagens e potencialidades, bem como a segurança dos seus utilizadores, ou destinatários das suas decisões²⁸⁸.

2. A necessidade histórica do regime de responsabilidade objetiva

Nos últimos dois séculos, a humanidade assistiu a uma sequência de novas descobertas, e desenvolvimentos científicos, que alteraram a nossa forma de organizar a sociedade. Pense-se no impacto do desenvolvimento da produção mecanizada, que transformou os meios de produção, no século XIX. Aquilo que, à época,

²⁸⁸ OLIVEIRA, ARLINDO; FIGUEIREDO, MÁRIO (2023). *Artificial Intelligence: Historical Context and State of the Art*. In “Multidisciplinary Perspectives on Artificial Intelligence and the Law”, eds. HENRIQUE SOUSA ANTUNES, PEDRO FREITAS, ARLINDO OLIVEIRA, CLARA MARTINS PEREIRA, ELSA VAZ DE SEQUEIRA & LUÍS BARRETO XAVIER, Springer, Law, Governance and Technology Series, vol. 58, pp. 3-24.

poderia parecer uma mera facilitação na produção fabril, tornou-se um verdadeiro impulsionador de um novo estilo de vida, e de trabalho.

No rescaldo da primeira Revolução Industrial, criou-se a consciência da necessidade de intervenção do Direito no funcionamento da sociedade. De facto, a evolução científica verificada desde a Revolução Industrial tem criado uma variedade de perigos, que podem pôr em risco a segurança dos cidadãos. No sentido de assegurar a sua proteção, foram estabelecidos, por um lado, deveres de cuidado imputáveis aos controladores de fontes de perigo. Por outro lado, no sentido de permitir o desenvolvimento de atividades económicas indispensáveis, mas as quais criavam riscos não domináveis, foram introduzidos mecanismos de imputação objetiva ao lesante, os quais seriam independentes da sua culpa.

A ideia histórica do estabelecimento destes mecanismos consistiria em garantir que, embora permitindo o desenvolvimento de atividades com perigos inerentes, tais atividades não afetariam excessivamente os cidadãos sujeitos a tal risco, aliviando assim o

ónus da prova, para a determinação da responsabilidade civil do controlador da atividade perigosa.

Tal foi a *ratio* subjacente à aprovação dos regimes de responsabilidade objetiva, tanto no regime dos acidentes de viação (artigo 503.º, Código Civil), como no regime da responsabilidade do produtor (DL n.º 383/89, de 06 de novembro).

No regime dos acidentes de viação, a responsabilidade objetiva foi consagrada pelo Decreto n.º 5.646, de 10 de maio de 1919, sendo este regime de imputação objetiva justificado pela maior intensidade do tráfego rodoviário, o qual era acompanhado pela maior frequência de falhas mecânicas, e de acidentes. A ideia subjacente ao regime da responsabilidade objetiva dos acidentes de viação consiste no facto de os veículos automóveis apresentarem um grau de perigo significativo, pela gravidade dos danos que possam causar²⁸⁹. Sen-

²⁸⁹ ATAÍDE, RUI; RODRIGUES, ANTÓNIO (2022). *Acidentes de Viação. Responsabilidade Subjetiva, Presunções de Culpa e Responsabilidade Objetiva*. In Revista Julgar, Almedina, vol. 46, pp. 13-32.

do este risco permitido, apresenta maiores perigos para os outros cidadãos, pelo que será contrabalançado com um regime simplificado de responsabilidade civil. Simultaneamente, a opção pelo regime de responsabilidade objetiva estará também relacionada com a frequente dificuldade de identificação da concreta causa do acidente, o que dificultaria a prova da culpa.

Embora tal limitação pudesse ser colmatada pela aplicação de um regime de responsabilidade subjetiva, com presunção de culpa, o que se pretende é que o detentor do veículo responda, não apenas pelos danos causados diretamente pela sua conduta, como também pelos danos derivados da própria posse e utilização do veículo²⁹⁰. No fundo, a imputação objetiva é justificada pela escolha em possuir e deter um objeto potencialmente perigoso, sendo que os danos resultantes de tal escolha não deverão ser suportados pelos restantes

concidadãos, caso a sua conduta não seja condenável.

Relativamente ao regime de responsabilidade objetiva do produtor (DL n.º 383/89, de 6 de novembro, que transpõe a Diretiva 85/374/CEE, em matéria de responsabilidade decorrente de produtos defeituosos), este pretende compensar o desequilíbrio negocial entre produtor, e consumidor, quanto à informação por estes detida. Consciente do risco de desproteção do consumidor, por desconhecimento do defeito de um produto, a União Europeia visou facilitar ao mesmo a compensação, por danos causados pelo mesmo²⁹¹. O equilíbrio é obtido através da possibilidade, dada ao produtor, de afastar a sua responsabilidade, caso demonstre, nomeadamente, que não colocou o produto em circulação, que este não possuía o defeito no momento em que foi colocado em circulação, ou que tal defei-

²⁹¹ COELHO, VERA (2017). *Responsabilidade do produtor por produtos defeituosos, “Teste de resistência” ao DL n.º 383/89, de 6 de novembro, à luz da jurisprudência recente, 25 anos volvidos sobre a sua entrada em vigor*. In Revista Electrónica de Direito, Faculdade de Direito da Universidade do Porto, vol. 2.

²⁹⁰ ATAÍDE, RUI; RODRIGUES, ANTÓNIO (2022). *Acidentes de Viação. Responsabilidade Subjetiva, Presunções de Culpa e Responsabilidade Objetiva*. In Revista Julgar, Almedina, vol. 46, pp. 13-32.

to não seria detetável, à luz dos conhecimentos científicos e técnicos, no momento em que foi colocado a circular²⁹².

Deste modo, a UE pretendeu garantir que se mantinham os incentivos económicos ao desenvolvimento de atividades industriais, mas que os consumidores não seriam prejudicados pelo desequilíbrio negocial natural da sua posição. Simultaneamente, ao facilitar a defesa dos consumidores, fomenta-se uma maior confiança dos mesmos, a qual também é benéfica ao desenvolvimento do comércio.

De modo semelhante, com o desenvolvimento da tecnologia de Inteligência Artificial, criou-se um verdadeiro desequilíbrio entre as partes. De facto, os utilizadores desta tecnologia raramente compreendem o seu funcionamento, tornando-se assim difícil provar os pressupostos da responsabilidade civil subjetiva.

²⁹² COELHO, VERA (2017). *Responsabilidade do produtor por produtos defeituosos, “Teste de resistência” ao DL n.º 383/89, de 6 de novembro, à luz da jurisprudência recente, 25 anos volvidos sobre a sua entrada em vigor*. In Revista Electrónica de Direito, Faculdade de Direito da Universidade do Porto, vol. 2.

Seguidamente, serão analisadas as fontes de risco dos sistemas de IA, que justificam uma aplicação do regime de responsabilidade objetiva, relativamente a certos sistemas de IA.

3. As características particulares da Inteligência Artificial

3.1. A falta de explicabilidade

A mais significativa fonte de desequilíbrio, e consequentemente, de risco, resulta da disparidade de informação do utilizador, em relação ao produtor do sistema de Inteligência Artificial. A desigualdade relativa ao conhecimento do modo de funcionamento de um produto ou serviço torna-se um fenómeno cada vez mais comum, dada a maior complexidade dos produtos atuais²⁹³.

²⁹³ NOGAROLI, RAFAELLA; FALEIROS JR., JOSÉ (2023). *Ethical Challenges of Artificial Intelligence in Medicine and the Triple Semantic Dimensions of Algorithmic Opacity with Its Repercussions to Patient Consent and Medical Liability*. In “Multidisciplinary Perspectives on Artificial Intelligence and the Law”, eds. HENRIQUE SOUSA ANTUNES, PEDRO FREITAS, ARLINDO OLIVEIRA, CLARA MARTINS PEREIRA, ELSA VAZ DE SEQUEIRA & LUÍS BARRETO XAVIER, Springer,

De facto, a maioria dos utilizadores não saberá, pelo menos na totalidade, como funcionam os produtos eletrónicos que adquire, ou os próprios componentes do seu carro. Isto resulta num maior controlo e poder, por parte do produtor, que poderá mais facilmente ocultar defeitos ou erros do seu produto. Tal levará, por um lado, a que o consumidor que os adquira não faça uma escolha informada, e por outro, irá dificultar a obtenção de compensação por parte do utilizador, visto que lhe será mais difícil explicar qual a circunstância que terá levado ao dano, bem como afastar uma potencial utilização incorreta do produto, pela sua parte.

Este desequilíbrio de posições acentua-se quanto a produtos ou serviços que utilizem tecnologias de Inteligência Artificial. Os avanços no desenvolvimento da Inteligência Artificial vêm a torná-la numa tecnologia cada vez mais complexa, e por essa razão, mais dificilmente compreensível para o utilizador comum²⁹⁴. Si-

Law, Governance and Technology Series, vol. 58, pp. 229-250.

²⁹⁴ “The behaviour of complex systems is difficult for us to grasp and often appears counter-intuitive. It is exempli-

multaneamente, os sistemas mais complexos de IA têm vindo a revelar-se caracterizados por uma opacidade, ou falta de transparência, quanto ao seu funcionamento²⁹⁵.

Em particular, os sistemas de *deep learning*²⁹⁶, que utilizam

fied by the famous butterfly effect, where the sensitive dependence on initial conditions can result in large differences at a later stage, as when the flapping of a butterfly’s wings in the Amazon leads to a tornado making landfall in Texas”, NOWOTNY, HELGA (2021). In *AI We Trust: Power, Illusion and Control of Predictive Algorithms*. Polity, p. 4. | OLIVEIRA, ARLINDO (2019). *Mentes Digitais: A Ciência Redefinindo a Humanidade*. (3.ª edição). IST Press, p. 83.

²⁹⁵ GONÇALVES-SÁ, JOANA; PINHEIRO, FLÁVIO (2023). *Societal Implications of Recommendation Systems: A Technical Perspective*. In “Multidisciplinary Perspectives on Artificial Intelligence and the Law”, eds. HENRIQUE SOUSA ANTUNES, PEDRO FREITAS, ARLINDO OLIVEIRA, CLARA MARTINS PEREIRA, ELSA VAZ DE SEQUEIRA & LUÍS BARRETO XAVIER, Springer, Law, Governance and Technology Series, vol. 58, pp. 47-64.

²⁹⁶ OLIVEIRA, ARLINDO; FIGUEIREDO, MÁRIO (2023). *Artificial Intelligence: Historical Context and State of the Art*. In “Multidisciplinary Perspectives on Artificial Intelligence and the Law”, eds. HENRIQUE SOUSA ANTUNES, PEDRO FREITAS, ARLINDO OLIVEIRA, CLARA MARTINS PEREIRA, ELSA VAZ DE SEQUEIRA & LUÍS BARRETO XAVIER, Springer,

múltiplas camadas de redes neurais, procurando assemelhar o seu funcionamento ao de neurónios humanos, têm apresentado resultados imprevisíveis, ou mesmo incompreensíveis. Este efeito é denominado de *black-box*, o qual pretende descrever a falta de transparência de certos sistemas de Inteligência Artificial. Esta falta de transparência concretiza-se no facto de que, embora sejam conhecidos os dados inseridos no sistema (*inputs*), bem como os resultados por este produzidos (*outputs*), são frequentemente desconhecidos os critérios por este utilizados, para obter um certo resultado²⁹⁷.

A falta de transparência não constituiria um obstáculo tão significativo, se os sistemas de Inteligência Artificial revelassem uma precisão infalível, demonstrando não cometer quaisquer erros ou discriminações. No entanto, sistemas de Inteligência Artificial já produziram, múlti-

Law, Governance and Technology Series, vol. 58, pp. 3-24.

²⁹⁷ BROZEK, BARTOS; JAKUBIEC, MAREK, KUCHARZYK, BARTŁOMIEJ (2024). *The Black Box Problem Revisited: Real and Imaginary Challenges for Automated Legal Decision Making*. In *Artificial Intelligence and Law*, vol. 32, pp. 427-440.

plas vezes, resultados erróneos ou discriminatórios.

A título de exemplo, refira-se uma experiência em que um sistema de Inteligência Artificial pretendia categorizar fotografias de animais, distinguindo entre cães e lobos. Inicialmente, o sistema parecia produzir resultados fidedignos, mas a certo ponto começou a apresentar resultados incorretos, categorizando erroneamente as duas espécies de animais. Os cientistas aperceberam-se, então, que o critério utilizado pelo sistema para a distinção não se referia às características do animal em si, mas às do cenário envolvente. Deste modo, o sistema caracterizava como lobos os animais que se encontravam rodeados por neve, e como cães os restantes animais²⁹⁸. Este exemplo pretende demonstrar como a falta de transparência dos sistemas poderá ocultar a utilização pelos mesmos de critérios inadequados para a tomada de uma decisão.

²⁹⁸ Esta experiência foi realizada no UCI Department of Computer Science, e os seus detalhes podem ser consultados em: <https://innovation.uci.edu/2017/08/husky-or-wolf-using-a-black-box-learning-model-to-avoid-adoption-errors/>

Um exemplo mais impactante poderia ser encontrado num sistema de diagnóstico de risco de pacientes com pneumonia. O sistema concluiu que os pacientes com asma tinham um risco baixo, quando contraíam pneumonia, pelo mero facto de recuperarem mais rapidamente, após a doença. Deste modo, o sistema de IA ignorou que os pacientes com asma, geralmente, recebem mais cuidados, o que facilita a respetiva recuperação²⁹⁹.

Através destes exemplos, pretende ilustrar-se como a falta de transparência, devida ao efeito *black-box*, consiste numa das características da Inteligência Artificial que acentua o desequilíbrio entre o utilizador e o produtor, sustentando-se por isso, a existência de um risco significativo, que justifica a aplicação de um regime de responsabilidade objetiva.

Seguidamente, será analisada a questão dos enviesamentos e discriminação, enquanto fator de risco adicional, inerente aos sistemas de Inteligência Artificial.

²⁹⁹ LAWRY, TOMROZEK, MUTKOSKI, STEVE, LEONG, NATHAN (2018) *Realizing the potential for AI in precision health*. In *SciTech Lawyer*, vol 15, pp. 23–27

3.2. Enviesamentos e discriminação

A atribuição de decisões a uma máquina criaria a convicção de que os resultados obtidos por esta serão caracterizados por uma maior neutralidade. De facto, encontramos-nos habituados a suspeitar de vieses em decisões humanas, mas tal ideia parecerá mais improvável quanto a decisões tomadas autonomamente por um sistema de IA.

Porém, como já foi referido, a realidade tem demonstrado que os sistemas de IA são capazes de produzir decisões discriminatórias. Tal discriminação poderá ocorrer por uma de duas fontes: ou na programação do sistema, ou por via da própria aprendizagem do sistema de Inteligência Artificial³⁰⁰.

³⁰⁰ MAGRANI, EDUARDO; SILVA, PAULA (2023). *The Ethical and Legal Challenges of Recommender Systems Driven by Artificial Intelligence*. In “Multidisciplinary Perspectives on Artificial Intelligence and the Law”, eds. HENRIQUE SOUSA ANTUNES, PEDRO FREITAS, ARLINDO OLIVEIRA, CLARA MARTINS PEREIRA, ELSA VAZ DE SEQUEIRA & LUÍS BARRETO XAVIER, Springer, Law, Governance and Technology Series, vol. 58, pp. 141-168.

Relativamente à primeira fonte de discriminação, será apenas relevante salientar a influência que a criação de um sistema inerentemente discriminatório poderá ter. O facto de ser possível ocultar uma decisão conscientemente discriminatória, através da decisão aparentemente neutra de um sistema de IA, poderá aumentar as situações de discriminação, deixando simultaneamente os destinatários das decisões indefesos, pela dificuldade inerente à imputação da decisão.

Quanto à segunda fonte de enviesamento, esta resulta da forma como se processa o treinamento de sistemas de IA. De facto, o treino dos sistemas baseia-se na alimentação do sistema com dados, a partir dos quais estes retirarão os critérios a utilizar nas decisões futuras³⁰¹. Porém, mesmo que o algoritmo não apre-

³⁰¹ OLIVEIRA, ARLINDO; FIGUEIREDO, MÁRIO (2023). *Artificial Intelligence: Historical Context and State of the Art*. In “Multidisciplinary Perspectives on Artificial Intelligence and the Law”, eds. HENRIQUE SOUSA ANTUNES, PEDRO FREITAS, ARLINDO OLIVEIRA, CLARA MARTINS PEREIRA, ELSA VAZ DE SEQUEIRA & LUÍS BARRETO XAVIER, Springer, Law, Governance and Technology Series, vol. 58, pp. 3-24.

sente enviesamentos, os dados utilizados para o treinar (*training set*) revelam frequentemente os preconceitos existentes na sociedade, principalmente no que se refere a características como a idade, o sexo, a raça, ou fatores económicos ou sociais.³⁰² Assim, frequentemente, os conjuntos de dados com os quais o sistema é alimentado, revelam uma distribuição de características, que não corresponde à distribuição existente na sociedade³⁰³.

³⁰² FREITAS, ANA TERESA (2023). *Data-Driven Approaches in Healthcare: Challenges and Emerging Trends*. In “Multidisciplinary Perspectives on Artificial Intelligence and the Law”, eds. HENRIQUE SOUSA ANTUNES, PEDRO FREITAS, ARLINDO OLIVEIRA, CLARA MARTINS PEREIRA, ELSA VAZ DE SEQUEIRA & LUÍS BARRETO XAVIER, Springer, Law, Governance and Technology Series, vol. 58, pp. 65-80.

³⁰³ GONÇALVES-SÁ, JOANA; PINHEIRO, FLÁVIO (2023). *Societal Implications of Recommendation Systems: A Technical Perspective*. In “Multidisciplinary Perspectives on Artificial Intelligence and the Law”, eds. HENRIQUE SOUSA ANTUNES, PEDRO FREITAS, ARLINDO OLIVEIRA, CLARA MARTINS PEREIRA, ELSA VAZ DE SEQUEIRA & LUÍS BARRETO XAVIER, Springer, Law, Governance and Technology Series, vol. 58, pp. 47-64.

Além disso, os dados poderão revelar algum tipo de preconceitos sociais, que serão assim expandidos, e propagados pelos sistemas de IA. Refira-se ainda que a seleção do conjunto dos dados de treino implica uma escolha, por parte dos programadores, que possivelmente irá refletir o seu sistema de valores e crenças³⁰⁴.

Todos estes fatores irão potenciar o enviesamento das decisões tomadas pelo sistema de IA. A discriminação resultante de decisões tomadas por sistemas de IA será particularmente grave, dada as características dos mesmos. Em primeiro lugar, foi já referida a falta de transparência, quanto aos critérios utilizados por um sistema, na tomada de uma decisão, o que dificulta a perceção quanto à utilização de critérios discriminatórios, proibidos por lei. Adicionalmente,

³⁰⁴ GONÇALVES-SÁ, JOANA; PINHEIRO, FLÁVIO (2023). *Societal Implications of Recommendation Systems: A Technical Perspective*. In “Multidisciplinary Perspectives on Artificial Intelligence and the Law”, eds. HENRIQUE SOUSA ANTUNES, PEDRO FREITAS, ARLINDO OLIVEIRA, CLARA MARTINS PEREIRA, ELSA VAZ DE SEQUEIRA & LUÍS BARRETO XAVIER, Springer, Law, Governance and Technology Series, vol. 58, pp. 47-64.

mesmo que seja identificado o critério discriminatório, existirá um obstáculo ao nível da imputação da conduta que levou à decisão enviesada.

Assim, por um lado, será difícil ao lesado pela decisão obter compensação, e por outro lado, será facilitada a ocultação de preconceitos, na tomada de decisões automatizadas.

4. O regime de responsabilidade civil por danos causados por sistemas de IA

4.1. Enquadramento

A análise dos fatores de risco associados às características intrínsecas de certos sistemas de IA pretendeu sustentar a necessidade de estabelecimento de um regime de responsabilidade objetiva, para certos sistemas de IA.

De facto, a falta de explicabilidade e a tomada de decisões enviesadas, e discriminatórias, poderá produzir danos significativos nos destinatários das decisões dos sistemas de IA. Simultaneamente, o facto de as decisões não serem compreensíveis, irá colocar o lesado numa posição de significativa desvantagem, dado que lhe será difícil demonstrar que a decisão foi errada, ou

discriminatória.³⁰⁵ Além disso, este terá ainda de provar o nexo de causalidade entre a conduta ilícita de um dos intervenientes na produção, e programação do sistema, e o concreto dano. Entende-se a dificuldade de prova deste nexo de causalidade, para um cidadão comum³⁰⁶.

A *ratio* que justificou o estabelecimento de um regime de responsabilidade objetiva, para os acidentes de viação, e para os danos provocados por produtos, fundamentaria a aplicação de um

³⁰⁵ MAGRANI, EDUARDO; SILVA, PAULA (2023). *The Ethical and Legal Challenges of Recommender Systems Driven by Artificial Intelligence*. In “Multidisciplinary Perspectives on Artificial Intelligence and the Law”, eds. HENRIQUE SOUSA ANTUNES, PEDRO FREITAS, ARLINDO OLIVEIRA, CLARA MARTINS PEREIRA, ELSA VAZ DE SEQUEIRA & LUÍS BARRETO XAVIER, Springer, Law, Governance and Technology Series, vol. 58, pp. 141-168.

³⁰⁶ FONSECA, ANA TAVEIRA; VAZ DE SEQUEIRA, ELSA; XAVIER, LUÍS BARRETO (2023). *Liability for AI Driven Systems*. In “Multidisciplinary Perspectives on Artificial Intelligence and the Law”, eds. HENRIQUE SOUSA ANTUNES, PEDRO FREITAS, ARLINDO OLIVEIRA, CLARA MARTINS PEREIRA, ELSA VAZ DE SEQUEIRA & LUÍS BARRETO XAVIER, Springer, Law, Governance and Technology Series, vol. 58, pp. 299-317.

regime de responsabilidade objetiva, para os sistemas de IA que revelassem um perigo acrescido. De facto, o risco associado aos veículos automóveis será equiparado ao risco dos danos que um sistema de IA poderá provocar, não só ao nível da própria integridade física da pessoa, por exemplo, quando a sua utilização se verifique no setor da saúde ou dos carros autónomos, como também a nível do potencial impacto dos sistemas em decisões relativas à admissão a uma universidade, a um emprego, ou na concessão de crédito, as quais poderão ser verdadeiramente definidoras na vida do destinatário da decisão³⁰⁷.

³⁰⁷ GONÇALVES-SÁ, JOANA; PINHEIRO, FLÁVIO (2023). *Societal Implications of Recommendation Systems: A Technical Perspective*. In “Multidisciplinary Perspectives on Artificial Intelligence and the Law”, eds. HENRIQUE SOUSA ANTUNES, PEDRO FREITAS, ARLINDO OLIVEIRA, CLARA MARTINS PEREIRA, ELSA VAZ DE SEQUEIRA & LUÍS BARRETO XAVIER, Springer, Law, Governance and Technology Series, vol. 58, pp. 47-64. | SOUSA ANTUNES, HENRIQUE (2021). *A Responsabilidade Civil aplicável à Inteligência Artificial: Primeiras Notas Críticas sobre a Resolução do Parlamento Europeu de 2020*. In Revista de Direito da Responsabilidade, Ano 3.

Por outro lado, o risco associado aos danos provocados por produtos aos seus consumidores prende-se, não apenas com o perigo da existência de defeitos no produto, mas principalmente com o desequilíbrio de informação entre o produtor, e o consumidor. Aliás, refira-se que a potencial classificação dos sistemas de IA como produtos possibilitaria, inclusive, com algumas necessárias adaptações, a inclusão dos mesmos no âmbito de aplicação da Diretiva da Responsabilidade do Produtor.

Dado o impacto que os sistemas de IA poderão ter na vida dos cidadãos, e a simultânea posição de desvantagem em que estes estarão, para poder obter compensação pelos danos provocados pelos sistemas, sustenta-se a existência de um risco significativo, que justifica a aplicação de um regime de responsabilidade objetiva. Tal aplicação estaria sujeita à definição de um critério de risco, que tomasse em conta, por um lado, as características intrínsecas dos sistemas, com maior relevância, a sua falta de explicabilidade, e o risco de enviesamento, e por outro lado, a sua concreta área de aplicação.

Seguidamente, será criticamente analisado o regime concretamente proposto pela União Europeia, para a responsabilidade civil por danos provocados por sistemas de IA.

4.2. A Proposta da União Europeia (Análise da Resolução do Parlamento Europeu de 2020 e Diretiva do Parlamento Europeu e do Conselho relativa à adaptação das regras de responsabilidade civil extracontratual à inteligência artificial)

4.2.1. A distinção com base no risco

A Resolução do Parlamento Europeu, aprovada a 20 de outubro de 2020, relativa ao regime de responsabilidade civil aplicável à Inteligência Artificial consagrou algumas recomendações à Comissão, para a futura elaboração de uma Diretiva nesta matéria.

Já em setembro de 2022, tendo em consideração a prévia Resolução, foi adotada uma Proposta de Diretiva do Parlamento Europeu e do Conselho relativa à adaptação das regras de responsabilidade civil extracontratual à inteligência artificial.

Um dos aspetos relevantes, comum a ambos os documentos,

consiste na estipulação de normas especialmente aplicáveis aos sistemas de IA de alto risco. Os contornos específicos do regime proposto serão analisados no capítulo seguinte, cumprindo agora deter-nos sobre o conceito de sistema de IA de alto risco.

A Resolução do Parlamento Europeu definiu o conceito de “sistema de IA de alto risco”, artigo 3.º, alínea c), como sendo caracterizado pelo seu funcionamento autónomo, e pela aleatoriedade na forma como causaria prejuízos.

O artigo 2.º, n.º 2, da Proposta de Diretiva utiliza uma denominação diversa, distinguindo os sistemas de IA de risco elevado, mas remetendo para o artigo 6.º do Regulamento da Inteligência Artificial (*AI Act*), o qual define o conceito de sistema de IA de alto risco, pelo que será seguro presumir a sinonímia entre as duas denominações.

Para a classificação de sistemas de IA como sistemas de alto-risco, o Regulamento de Inteligência Artificial utiliza dois critérios alternativos.

O primeiro critério estabelece a classificação, como sistemas de IA alto risco, dos sistemas que preencherem duas condições

cumulativas: a previsão do produto na listagem do Anexo I ao Regulamento de Inteligência Artificial, e a sujeição do sistema à verificação de conformidade, por parte de terceiros (artigo 6.º n.º 1, alíneas a) e b)). Refira-se que todas as utilizações mencionadas no Anexo I se baseiam na área da aplicação concreta do sistema de IA, sendo exemplos de tal a utilização de sistemas de IA no campo da aviação civil, nos veículos, ou em instalações por cabo.

Enquadrados no segundo critério, são categorizados como sistemas de alto risco os previstos no Anexo III, se comportarem num risco significativo de dano para a saúde, segurança ou para os direitos fundamentais (artigo 6.º n.ºs 2 e 3). A título de exemplo, refira-se que se consideram como sistemas de risco elevado aqueles que realizam categorização biométrica, de acordo com atributos ou características sensíveis, Anexo III, n.º 1, alínea b), assim como os sistemas de reconhecimento de emoções, Anexo III, n.º 1, alínea c). Assim, compreende-se que a classificação do Anexo III pretende incluir as utilizações de sistemas que incluam tecnologias de Inteligência Artificial que lidem com áreas de maior sensi-

bilidade. Para compensar o risco acrescido destes sistemas, serão os produtores dos mesmos a ter de comprovar que, ainda que o seu sistema se inclua na previsão do Anexo III, este não representará um risco significativo para os bens supramencionados.

A categorização enquanto sistema de IA de alto risco releva, no sentido em que comporta obrigações adicionais para o seu produtor, à luz do Regulamento de Inteligência Artificial. Com maior destaque, refira-se a obrigação de transparência (artigo 13.º, Regulamento de Inteligência Artificial), de supervisão humana (artigo 14.º, Regulamento de Inteligência Artificial), e de cibersegurança do sistema (artigo 15.º, Regulamento de Inteligência Artificial).

É relevante notar que, enquanto a Resolução colocava a tónica da definição no modo de funcionamento do sistema, o Regulamento de IA veio focar-se no setor de utilização do sistema, e na concreta natureza das atividades³⁰⁸. Ainda que a Resolução

também instasse a tomar em conta estes fatores, o seu foco nas características internas do sistema comportaria uma visão mais adequada, na minha opinião, para a definição de existência de um risco. E tal distinção seria, em particular, relevante para a definição de um regime diferenciado de responsabilidade civil, como será de seguida sustentado.

4.2.2. O regime de responsabilidade objetiva para os sistemas de IA de risco elevado

Na Resolução, diferentemente do que sucede na Proposta de Diretiva, propõe-se um regime de responsabilidade objetiva para os sistemas de IA de risco elevado³⁰⁹.

Ambos os documentos reconhecem o facto de que a opacidade dos sistemas tornaria consideravelmente difícil, ou mesmo impossível, identificar o autor

Críticas sobre a Resolução do Parlamento Europeu de 2020. In Revista de Direito da Responsabilidade, Ano 3.

³⁰⁹ SOUSA ANTUNES, HENRIQUE (2021). *A Responsabilidade Civil aplicável à Inteligência Artificial: Primeiras Notas Críticas sobre a Resolução do Parlamento Europeu de 2020. In Revista de Direito da Responsabilidade, Ano 3.*

da conduta, que veio a provocar o dano (considerando 7, Resolução), considerando as características de certos sistemas de IA, em particular, a falta de transparência, resultante numa difícil explicabilidade das suas decisões. No entanto, a mesma premissa fundamentou propostas diferentes, como se irá verificar.

Relativamente ao regime proposto na Resolução, será criticável o seu efeito, quando conjugado com a definição de sistema de IA de alto risco, do Regulamento de IA. O facto de ser utilizado um conceito de sistema de IA de alto risco, maioritariamente baseado na área de aplicação dos sistemas, determina que não são verdadeiramente consideradas as características intrínsecas dos sistemas de IA, na determinação da aplicação de um regime de responsabilidade objetiva. Tal lacuna na definição desconsidera os sistemas de IA que, ainda que não se insiram na lista do Anexo III, e que não sejam aplicados em áreas de maior sensibilidade, possam comportar um risco significativo, pelas suas próprias características intrínsecas. Em particular, seria relevante incluir na definição de sistema de IA de alto risco os sistemas que reve-

lem uma opacidade quanto aos critérios utilizados na tomada de decisões, dado que esse é inclusive um fator de risco reconhecido, tanto na Resolução, como na Proposta de Diretiva.

Além disso, à luz da Resolução, seriam também sujeitos a responsabilidade objetiva os sistemas de IA que não tivessem sido avaliados pela Comissão, e pelo Comité, e que, por tal razão, não constassem ainda como sistemas de alto risco, mas que, em incidentes repetidos, houvessem já provocado danos ou prejuízos graves.

Refira-se ainda que a responsabilidade objetiva seria afastada, nos casos em que os prejuízos ou danos fossem causados por motivos de força maior (artigo 4.º, n.º 3, Resolução), o que se justifica, dado que, pela sua imprevisibilidade, e pela incontroabilidade do produtor sobre os mesmos, não deverá ser um risco por este suportado.

4.2.3. O regime de responsabilidade culposa com presunção de culpa para os restantes sistemas de IA

O Parlamento propôs, em relação aos sistemas que não sejam

³⁰⁸ SOUSA ANTUNES, HENRIQUE (2021). *A Responsabilidade Civil aplicável à Inteligência Artificial: Primeiras Notas*

categorizados como sistemas de IA de risco elevado um regime de responsabilidade culposa, com presunção de culpa do produtor. Embora remeta a generalidade do regime aplicável para os regimes nacionais, nomeadamente quanto preenchimento dos pressupostos concretos, a Resolução estabelece um desvio ao regime geral de responsabilidade culposa.

De facto, na Resolução propõe-se uma presunção de culpa, por parte do operador, a qual poderá ser ilidida, através da prova do cumprimento de deveres de diligência (artigo 8.º n.º 2, Resolução)³¹⁰. Em particular, o operador poderá demonstrar que o sistema de IA foi ativado sem o seu conhecimento, provando, no entanto, que foram tomadas medidas para evitar tal ativação, ou poderá ainda demonstrar que terá sido diligente, na seleção do sistema

de IA, na sua correta colocação e operação, e no controlo e manutenção do mesmo. Seria relevante ponderar se o Parlamento pretendia consagrar uma lista taxativa quanto à possibilidade de afastar esta presunção de culpa. A redação da norma parece sugerir que as únicas possibilidades de afastar a presunção de culpa serão as expressamente previstas. Porém, tal tornar-se-ia verdadeiramente restritivo para o produtor de Inteligência Artificial, dado que o mesmo se encontra sujeito a uma presunção de culpa.

Além disso, poderá parecer excessivo que a presunção de culpa só possa ser afastada pela verificação da mais completa diligência em todas as fases do processo de instalação do sistema, e inclusive após a mesma. *A contrario*, poderíamos concluir que a falta de observância de diligência em um só destes momentos impediria o afastamento da presunção de culpa. Novamente, tal previsão parece realizar uma distribuição desproporcional de ónus entre o produtor e o utilizador.

A Resolução prevê ainda um regime de responsabilidade subsidiária do produtor, mesmo no caso em que se prove que o prejuízo foi causado pela interven-

ção de um terceiro, que haja alterado o funcionamento do sistema, e os seus efeitos, se o mesmo não for localizável ou carecer de recursos financeiros (artigo 8.º, n.º 3, Resolução). Esta previsão é altamente criticável, visto parecer estabelecer um regime de responsabilidade objetiva, subsidiária, sem ser possível justificar qual a fonte do risco³¹¹.

De facto, esta previsão só seria aplicável a sistemas que não hajam sido classificados como sistemas de IA de risco elevado, estando, portanto, sujeitos ao regime de responsabilidade culposa, o qual já inclui uma presunção de culpa. Sendo que o operador já terá afastado a sua responsabilidade, ao demonstrar que o dano se deveu à intervenção de um terceiro, que terá alterado o funcionamento e os efeitos do sistema, tornando-os, portanto, imprevisí-

veis para o operador, não se justifica que este continue a ter de compensar o lesado.

Não poderíamos classificar esta responsabilidade como culposa, dado ter sido afastada a presunção de culpa, e visto que não terá ficado provada qualquer conduta ilícita do operador que tenha desencadeado o dano. E também não seria aqui justificável a aplicação de um regime de responsabilidade objetiva. Na realidade, tal regime não é aplicado, nas circunstâncias em que o dano é provocado pelo mesmo sistema, sem a intervenção de terceiro, por se considerar que o risco que o sistema comporta não é considerável. A aplicação do regime de responsabilidade objetiva ao caso de intervenção de um terceiro levar-nos-ia a concluir pela existência de um risco na própria intervenção do terceiro. Ora, se o operador demonstrou o cumprimento dos deveres de diligência, nomeadamente quanto ao controlo das atividades e à manutenção da fiabilidade do sistema, este não deverá ser responsabilizado pela intervenção de um terceiro num sistema a que, provavelmente, já não tem acesso. Esse risco correria já por conta do seu utilizador, cumprimen-

³¹⁰ FONSECA, ANA TAVEIRA; VAZ DE SEQUEIRA, ELSA; XAVIER, LUÍS BARRETO (2023). *Liability for AI Driven Systems*. In “Multidisciplinary Perspectives on Artificial Intelligence and the Law”, eds. HENRIQUE SOUSA ANTUNES, PEDRO FREITAS, ARLINDO OLIVEIRA, CLARA MARTINS PEREIRA, ELSA VAZ DE SEQUEIRA & LUÍS BARRETO XAVIER, Springer, Law, Governance and Technology Series, vol. 58, pp. 299-317.

³¹¹ FONSECA, ANA TAVEIRA; VAZ DE SEQUEIRA, ELSA; XAVIER, LUÍS BARRETO (2023). *Liability for AI Driven Systems*. In “Multidisciplinary Perspectives on Artificial Intelligence and the Law”, eds. HENRIQUE SOUSA ANTUNES, PEDRO FREITAS, ARLINDO OLIVEIRA, CLARA MARTINS PEREIRA, ELSA VAZ DE SEQUEIRA & LUÍS BARRETO XAVIER, Springer, Law, Governance and Technology Series, vol. 58, pp. 299-317.

do a este a vigilância do próprio sistema.

A Proposta de Diretiva não adotou este regime, o que parece sensato, dado que este representaria um fardo excessivo nos produtores de sistemas de IA, que nem são classificados como sistemas de alto risco. A aplicação desta previsão poderia comportar um impedimento significativo, e injustificado, no desenvolvimento dos sistemas de IA.

4.2.4. O regime de responsabilidade culposa e a presunção de nexos de causalidade da Proposta de Diretiva

A Proposta de Diretiva, contrariando o sentido apresentado na Resolução do Parlamento, propõe um regime de responsabilidade culposa, introduzindo como elementos inovadores a divulgação de elementos de prova, por ordem do tribunal, assim como uma presunção de nexos de causalidade.

A respeito do regime de responsabilidade culposa será relevante apontar que a sua aplicação não parece salvaguardar os direitos dos cidadãos. De facto, os riscos inerentes às tecnologias de IA, bem como o desequilíbrio

por estes acentuado, justificaria uma maior proteção dos seus utilizadores. A necessidade de provar a própria culpa dos produtores, considerando a complexidade estrutural da indústria que origina um sistema de IA, bem como o desconhecimento relativo ao próprio funcionamento do sistema, poderá significar uma desproteção das vítimas de danos causados por estes sistemas.

A Proposta de Diretiva estabelece apenas uma presunção de culpa, nos casos em que o demandado incumpra a ordem do tribunal de divulgar ou conservar elementos de prova de que disponha (artigo 3º n.º5, Proposta de Diretiva). No entanto, tal presunção não parece suficiente para acautelar a tutela do lesado, dado que a divulgação de informações não determinará, provavelmente, a efetiva compreensão do funcionamento do sistema, nem a deteção de algum defeito ou falha do mesmo.

Por outro lado, a Proposta de Diretiva vem introduzir um elemento não contemplado no regime de responsabilidade civil proposto pela Resolução do Parlamento, que se afigura promissor. Concretamente, estabelece uma presunção de nexos de causa-

lidade, nos casos em que o lesado prove a culpa do demandado, ou em que esta seja presumida (artigo 4.º, Proposta de Diretiva)³¹².

O estabelecimento desta presunção é sustentado pela consciência da dificuldade dos lesados em estabelecer um nexo de causalidade entre uma conduta concreta, e um dano. De facto, considerando o previamente mencionado efeito *black-box*, refletido na falta de opacidade e transparência de sistemas de IA mais complexos, tornar-se-á difícil, para o utilizador comum, a prova da relação entre uma conduta ilícita e um resultado danoso³¹³.

³¹² SOUSA ANTUNES, HENRIQUE (2021). *A Responsabilidade Civil aplicável à Inteligência Artificial: Primeiras Notas Críticas sobre a Resolução do Parlamento Europeu de 2020*. In *Revista de Direito da Responsabilidade*, Ano 3.

³¹³ FONSECA, ANA TAVEIRA; VAZ DE SEQUEIRA, ELSA; XAVIER, LUÍS BARRETO (2023). *Liability for AI Driven Systems*. In “Multidisciplinary Perspectives on Artificial Intelligence and the Law”, eds. HENRIQUE SOUSA ANTUNES, PEDRO FREITAS, ARLINDO OLIVEIRA, CLARA MARTINS PEREIRA, ELSA VAZ DE SEQUEIRA & LUÍS BARRETO XAVIER, Springer, Law, Governance and Technology Series, vol. 58, pp. 299-317.

Isto pois o processo decisório do sistema de IA poderá ser de tal modo complexo, que nem os seus programadores conhecem os critérios utilizados pelo sistema para produzir um certo resultado, a partir dos dados introduzidos. Em muitos casos, os sistemas de IA têm inclusive, produzido resultados incompreensíveis e imprevisíveis, como já foi exemplificado anteriormente. Deste modo, será ainda mais difícil que os utilizadores, que na sua maioria desconhecem o modo de funcionamento do sistema, consigam provar a correlação entre uma conduta ou omissão ilícita, e um concreto resultado danoso. O estabelecimento da presunção de nexos de causalidade tem como objetivo evitar que os utilizadores fiquem materialmente impedidos de obter compensação, devido à excessiva dificuldade em demonstrar os pressupostos da responsabilidade civil.

Como foi referido, a presunção de causalidade seria aplicável nos casos em que existisse culpa do demandado, ou em que o tribunal presumisse essa culpa, com base no incumprimento de uma decisão judicial de divulgação ou conservação de elementos de prova. Adicionalmente, ressalve-se que

só será estabelecida a presunção de causalidade caso se demonstre a probabilidade de que o facto culposos invocado haja influenciado o resultado danoso do sistema de IA, ou a inexistência do resultado relevante. Além disso, o lesado terá de provar que foi o sistema de IA a originar os danos³¹⁴. Verifica-se, assim, que não foi proposta uma inversão do ónus da prova, quanto ao nexo de causalidade, mas sim um alívio da prova deste pressuposto, que se revela proporcional, por equilibrar os interesses de ambas as partes.

Refira-se ainda que, quanto às ações intentadas contra fornecedores de sistemas de IA de alto risco, o lesado só beneficiará da presunção de causalidade, se provar o incumprimento de um dos deveres impostos, pelo Regulamento de IA, aos produtores

de sistemas de IA de alto risco (nomeadamente, a obrigação de transparência (artigo 13.º), de supervisão humana (artigo 14.º), ou de cibersegurança do sistema (artigo 15.º).

O considerando 22 da Proposta de Diretiva vem esclarecer que a presunção só se aplicará no caso de incumprimento dos deveres de diligência que se destinavam a proteger contra os danos ocorridos. Embora este esclarecimento limite a possibilidade de aplicações abusivas da presunção de culpa, será ainda difícil equacionar quais os deveres de diligência concretos que se destinavam a proteger contra um dano específico.

Na realidade, não se afigura acessível identificar qual será o dever cujo incumprimento estará diretamente correlacionado com um dano concreto, considerando a já referida falta de transparência quanto ao modo de funcionamento dos sistemas³¹⁵.

³¹⁴ FONSECA, ANA TAVEIRA; VAZ DE SEQUEIRA, ELSA; XAVIER, LUÍS BARRETO (2023). *Liability for AI Driven Systems*. In “Multidisciplinary Perspectives on Artificial Intelligence and the Law”, eds. HENRIQUE SOUSA ANTUNES, PEDRO FREITAS, ARLINDO OLIVEIRA, CLARA MARTINS PEREIRA, ELSA VAZ DE SEQUEIRA & LUÍS BARRETO XAVIER, Springer, Law, Governance and Technology Series, vol. 58, pp. 299-317.

³¹⁵ MAGRANI, EDUARDO; SILVA, PAULA (2023). *The Ethical and Legal Challenges of Recommender Systems Driven by Artificial Intelligence*. In “Multidisciplinary Perspectives on Artificial Intelligence and the Law”, eds. HENRIQUE SOUSA ANTUNES, PEDRO FREITAS,

Deste modo, ou o conteúdo da norma seria quase inaplicável, pela falta de verificação de um requisito que se concretizará, na prática, na difícil verificação do nexo de causalidade, ou teria como consequência a imposição de uma presunção excessiva e injustificada no produtor, pois o incumprimento de um simples dever de diligência poderá não estar remotamente relacionado com o dano concreto.

Visto que a Proposta de Diretiva vem estabelecer, no seu considerando 25, que a presunção só deverá ser aplicada quando se puder considerar, com razoável probabilidade, que o facto culposos influenciou o resultado, podemos concluir que os casos de aplicação desta norma seriam residuais.

Cumpra referir ainda que, no caso de sistemas de IA que não são de risco elevado, a presunção de causalidade só será estabelecida, caso o tribunal considere ser excessivamente difícil para o demandante provar o nexo de cau-

salidade. A Proposta de Diretiva esclarece que a dificuldade deverá ser apreciada à luz das características de determinados sistemas de IA, em particular, a sua autonomia e opacidade que, como foi supramencionado, irão dificultar a prova do nexo de causalidade (considerando 28).

4.2.5. O dever de prestar informações

Cumpra ainda referir, brevemente, a existência de um dever de cooperação do produtor do sistema de IA. Por decisão do juiz, este poderá estar sujeito a fornecer informações, que sejam relevantes para o pedido, no sentido de apurar a existência de responsabilidade (artigo 8.º, n.º 4, Resolução).

Por seu lado o artigo 3.º, n.º 1, da Proposta de Diretiva estabelece que pode ser ordenada pelo tribunal a divulgação de elementos de prova relevantes sobre sistemas de IA de risco elevado específicos, suspeitos de terem causado danos.

O estabelecimento de uma presunção de incumprimento dos deveres de diligência que esses elementos de prova se destinavam a provar, como consequência da

falta de divulgação dos mesmos (artigo 3.º, n.º 5, Proposta de Diretiva) visa incentivar a apresentação de provas. Tal norma, ainda que, como defendido, se revele insuficiente para colmatar as dificuldades de prova, por parte do lesado, é relevante, no sentido de garantir um maior acesso à informação, o qual, em última análise, trará um maior equilíbrio no âmbito do processo.

5. Conclusão

A definição de um regime de responsabilidade para os danos causados por sistemas de IA não se afigura simples, e será dificultada pela existência de uma multiplicidade de sistemas de IA, que comportam riscos diferenciados. Atingir um equilíbrio, entre a proteção dos utilizadores, e o incentivo ao desenvolvimento tecnológico, determina que sejam encontrados critérios para a definição do risco significativo, que delimite a aplicação de um regime de responsabilidade objetiva.

O Regulamento de Inteligência Artificial deu um primeiro passo nesse sentido, enumerando alguns sistemas de IA que seriam de alto risco, e elencando ainda

as áreas de aplicação que configurariam um risco superior. A Resolução do Parlamento vem propor a aplicação do regime de responsabilidade objetiva a estes sistemas.

Neste artigo, propõe-se o alargamento da responsabilidade objetiva a outros sistemas. De facto, com base nas consequências de algumas das características de sistemas IA apresentadas no decorrer do artigo, o critério definidor de sistema de IA de alto risco deveria considerar não apenas a área de aplicação do sistema, como também as suas características intrínsecas. Os sistemas de IA que apresentem uma falta de transparência, bem como tendência para enviesamentos, deverão ser considerados também como sistemas de IA de alto risco, para efeitos da aplicação de um regime de responsabilidade objetiva.

A combinação de um regime de responsabilidade objetiva, com o aligeiramento da prova do nexo de causalidade seria a solução defendida, que se afigura proporcional dados os elevados riscos, e a considerável magnitude dos danos potencialmente causados pela IA.

Só a aplicação deste regime poderá, por um lado, criar um

incentivo para que os produtores de sistemas de IA de alto risco assegurem a segurança, e fiabilidade dos seus sistemas, e por outro lado, garantir que os lesados pelas decisões dos mesmos têm uma possibilidade fáctica de obter compensação.

A procura de um regime equilibrado, e proporcional, de responsabilidade civil aplicada a danos provocados por sistemas de IA, irá em última análise, potenciar o desenvolvimento tecnológico, ao garantir uma maior confiança dos cidadãos na utilização desta tecnologia, a qual é fundamental para o incentivo ao progresso científico.