

Desafios internacionais da aplicação da inteligência artificial no Direito¹

International challenges of applying artificial intelligence in law

Felipe Chiarello²

Lara Rocha Garcia³

Resumo

O presente artigo analisa os impactos da inteligência artificial sem consciência na sociedade, em especial, no Direito. Com base em uma revisão de literatura do período de 2015 a 2020, são analisados os principais conceitos da inteligência artificial, na sequência, são examinados os documentos de diretrizes e regras para programação e desenvolvimento de softwares e, por fim, é avaliada a relação dos juristas com as máquinas em um cenário de *cyborg justice*.

Palavras-chave: *Cyborg Justice*. Direito Digital. Direito Econômico. Inteligência Artificial. Tecnologia.

¹ Recebido 10/1/2021. Aprovado: 21/07/2021.

² Possui mestrado e doutorado em Direito pela Pontifícia Universidade Católica de São Paulo. Foi Diretor da Faculdade de Direito e atualmente é Pró-Reitor de Pesquisa e Pós-Graduação da Universidade Presbiteriana Mackenzie, Professor Titular da Faculdade de Direito e do Programa de Mestrado e Doutorado em Direito Político e Econômico, membro da Academia Mackenzista de Letras, Professor Colaborador do Programa de Pós-graduação em Direito da Universidade de Passo Fundo (UPF), Coordenador Adjunto de Programas Acadêmicos da Área de Direito da CAPES-MEC e Membro Pesquisador 2 do CNPq. E-mail: felipe.chiarello@mackenzie.br.

³ Doutoranda e Mestre em Direito Político e Econômico pela Universidade Presbiteriana Mackenzie, Visiting Scholar pela Columbia Law School, com dupla formação em Comunicação Social pela Universidade Estadual Paulista “Júlio de Mesquita Filho” (UNESP) e Direito pela Universidade Presbiteriana Mackenzie; Especialista em Inovação por Stanford School of Business – Ignite Program e Visiting Scholar da Columbia Law School. Membro da Comissão de Direito Digital e Compliance da OAB/SP, foi Gerente do Laboratório de Inovação do Hospital Albert Einstein e Gerente de Produtos e Inovação do Dr.Consulta. E-mail: lararochagarcia@gmail.com.

Abstract

This article analyzes the impacts of unconscious artificial intelligence in society, especially in Law. Based on a literature review from the period 2015 to 2020, it analyses the main concepts of artificial intelligence, then the guidelines and rules documents for software programming and development are examined and, finally, it is the evaluated the relationship of jurists with the machines in a cyborg justice setting.

Keywords: Artificial Intelligence. Cyborg Justice. Digital Law. Economic Law. Technology.

Introdução

Richard Susskind, Mathias Risse, Klaus Schwab, David Chalmers e Jesica Fjeld são autores internacionais que se dedicaram, pelo menos nos últimos cinco anos, a estudar os impactos da inteligência artificial na sociedade, sob os aspectos econômicos, sociais, políticos e jurídicos.

Ainda que oscilem entre visões otimistas e pessimistas sobre este impacto, todos eles são convergentes em afirmar que a revolução homem-máquina mudará os rumos das relações humanas e tudo que dela decorre. O *como* ainda diverge, mas já existem 36 documentos escritos por governos, empresas, instituições e associações ao redor do mundo cuja tentativa é criar diretrizes para conduzir esse processo.

Este artigo pretende analisar a literatura internacional para entender o panorama da inteligência artificial desprovida de consciência e os desafios trazidos para as ciências jurídicas. Para isso, utilizou bibliografia internacional dos últimos cinco anos para analisar as diretrizes de programação de inteligência artificial, entender seus vieses, dificuldades e problemáticas gerais para, assim, ser possível compreender os impactos no Direito, especialmente no ramo do direito público.

Não pretende este artigo esgotar o tema, consciente de que ainda está em constante discussão, tem por objetivo trazer os conceitos principais e provocar uma discussão em âmbito nacional.

1. Potencialidades da Inteligência Artificial

Vivemos hoje a chamada quarta revolução industrial⁴, com a produção de uma economia pautada no digital e, como todas as revoluções que aconteceram, sejam agrícolas, industriais ou digitais, todas elas pautadas nas tecnologias emergentes de sua época, que proporcionou aumento de produtividade, eficiência e mudança comportamental.

Naturalmente, nessa revolução como nas outras, há a preocupação com o fim do trabalho como conhecemos, ou a sua transformação. Sobre isso Risse entende que a tecnologia acaba por criar mais trabalhos do que os destruir:

But so far every wave of technological innovation has ended up creating more jobs than it destroyed. While technological change was not good for everybody, it was good for society as a whole, and for humanity. It is possible that there will be so many jobs that those who develop, supervise or innovatively use technology, as well as creative professions that cannot be displaced, will eventually outnumber those who lose jobs to AI. ⁵

Assim, embora defenda que existem desafios para os indivíduos e que é preciso se preocupar com a igualdade (ou desigualdade) que pode ser proporcionada pela tecnologia, entende que, no geral, para o coletivo, a tecnologia é mais benéfica pois ela ajuda a desenvolver enquanto sociedade.

Richard Susskind tem dedicado os últimos trinta anos ao tema, desde sua primeira obra “Expert Systems in Law”⁶. Considerando o recorte metodológico desde artigo, que se trata do impacto de inteligência artificial na sociedade e, em especial, no Direito, avaliando as publicações dos últimos 5 anos (2015 a 2020), para este artigo, foram analisadas especialmente as obras *Online Courts and the Future of Justice*⁷, *Tomorrow’s Lawyers*⁸, segunda edição, e *The Future of the Professions*⁹.

⁴ SCHWAB, Klaus. **A quarta revolução industrial**. São Paulo: Edipro, 2018.

⁵ RISSE, Mathias; LIVINGSTON, Steven. The Future Impact of Artificial Intelligence on Humans and Human Rights. **Ethics and International Affairs**, n.33, v.2, p.141-158, jun. 2019. p.13

⁶ SUSSKIND, Richard. **Expert Systems in Law**. London: Oxford University Press, 1987.

⁷ SUSSKIND, Richard. **Online Courts and the Future of Justice**. London: Oxford University Press, 2019.

⁸ SUSSKIND, Richard. **Tomorrow’s Lawyers**. London: Oxford University Press, 2017.

⁹ SUSSKIND, Richard. **The Future of the Professions**. London: Oxford University Press, 2015.

Para Susskind, o futuro pode seguir dois caminhos, sendo o primeiro uma evolução incremental, sem grandes mudanças, mas com muito mais eficiência, produtividade e celeridade. No entanto, há um outro caminho, que gera mais transformação, que pondera outro papel para os juristas. O autor apresenta visões com o intuito de abrir os olhos para mudanças irreversíveis, como cortes judiciais virtuais, negócios globais inteiramente baseados em internet, produção de documentação online, serviços jurídicos *comoditizados* e simulação prática *web-based*. Para ele, teremos novos trabalhos para os advogados em mercados legais mais abertos, assim como novos tipos de profissionais não advogados nessa indústria.

Um dos responsáveis por essa mudança é a inteligência artificial:

AI is increasingly present in our lives, reflecting a growing tendency to turn for advice, or turn over decisions altogether, to algorithms. By “intelligence”, I mean the ability to make predictions about the future and solve complex tasks. “Artificial” intelligence, AI, is such ability demonstrated by machines, in smart phones, tablets, laptops, drones, self-operating vehicles or robots that might take on tasks ranging from household support, companionship of sorts, even sexual companionship, to policing and warfare.¹⁰

Impossível negar a existência de produtos, serviços e processos pautados em inteligência artificial, seria ingenuidade considerar que o Direito não seria impactado ou mesmo chamado a prover segurança jurídica.

No entanto, há que se entender a inteligência artificial minimamente para ser possível fazer análises mais profundas sobre seu futuro, afinal, como previu Souza e Oliveira, “Um mundo que oscila nas fronteiras entre o que aprendemos a tratar como ficção científica e um universo de possibilidades”¹¹

O primeiro ponto a se clarificar são as três dimensões técnicas¹²: presença de dados, capacidade de processamento e desenvolvimento de

¹⁰ RISSE, Mathias. Human Rights and Artificial Intelligence: An Urgently Needed Agenda. **Human Rights Quarterly**, n.41, v.1, p.1-16, fev. 2019. p.2.

¹¹ SOUZA, Carlos Affonso Pereira; OLIVEIRA, Jordan Vinicius. Sobre os ombros de robôs? A inteligência artificial entre fascínios e desilusões. In: FRAZÃO, Ana; MULHOLLAND, Caitlin. **Inteligência Artificial e Direito**. São Paulo: Thompson Reuters, 2019. p. 65

¹² STEIBEL, Fabro; VICENTE, Victor Freitas; JESUS, Diego Santos Vieira. Possibilidades e Potenciais da Utilização da Inteligência Artificial. In: FRAZÃO, Ana; MULHOLLAND, Caitlin. **Inteligência Artificial e Direito**. São Paulo: Thompson Reuters, 2019. p. 53.

softwares. A parte de hardware, ou seja, o substrato, não é necessário para ser considerado inteligência artificial e nem mesmo robô.

Portanto, resta claro a dependência da inteligência artificial com dados, sendo que o volume (*big data*), pode acelerar a programação. Por consequência, a capacidade de processamento da máquina está relacionada, como o próprio título sugere, ao tratamento destes dados de forma a concatená-los com resultados efetivos.

Importante diferenciar inteligência artificial de um sistema de árvore de decisão. Este último acontece no processo conhecido como if/then (se/então), que significa programar, em símbolos lógicos, por exemplo, a seguinte frase “se chover, então abra o guarda-chuva. Se não chover, então use a piscina”. Isso não é inteligência, é apenas lógica direta. Nesse caso, se ventar, por exemplo, o sistema não sabe o que fazer, ele simplesmente não fará nada porque, essa condição “se ventar” não está descrita no código. Se fosse um sistema inteligente, ele poderia fazer a analogia de vento com chuva, por exemplo, e abrir o guarda-chuva. Ainda que não seja uma atitude que um humano faria, ele possivelmente identificaria que isso se trata de um evento meteorológico e tentaria uma ação próxima a algo que ele conhece. Nesse caso, um humano poderia ensiná-lo que se trata de evento meteorológico diferente, que requer ação diferente. O que nos leva a entender os 3 tipos processo de aprendizado de máquina.

O aprendizado supervisionado significa um conjunto de inputs para um conjunto de resultados com métodos aplicados. Já o não supervisionado significa um *input* rotulado, mesmo que o resultado não seja, o que significa a necessidade de aferição. O terceiro, chamado de reforçado, trata-se de resultado variável para ser maximizado e uma série de decisões que podem ser tomadas.¹³

Estes aprendizados nos levam a entender as três necessidades, relacionadas com as capacidades técnicas já apresentadas, que seriam, primordialmente, a organização de dados, que significa o uso para estruturação;

¹³ STEIBEL, Fabro; VICENTE, Victor Freitas; JESUS, Diego Santos Vieira. Possibilidades e Potenciais da Utilização da Inteligência Artificial. p. 55.

o auxílio a tomada de decisão, que indica o uso de processamento e a automação da decisão, que demonstra a decisão em si.

Existem três graus de inteligência artificial¹⁴: a restrita, que, mesmo com esse nome, ainda concentra sistemas muito bons em alguma coisa, mas ruins em outras. Por exemplo, o Google, que é um buscador excepcional, mas não é um bom construtor de texto. Da mesma forma, o carro autônomo, que serve para movimentar um carro sem motorista, mas não serve para realizar uma cirurgia, ou identificar um melasma em uma base de imagens de tumores de pele.

O grau de inteligência chamado de “geral” se aproxima do que chamamos de consciente, ou o que as pessoas esperam de um software de inteligência artificial. Já o terceiro grau, o de “superinteligência” encontra-se em uma escala inumana, que pode ser considerada um risco para a própria humanidade.

Also, philosophers have long puzzled about the nature of the mind. One question is if there is more to the mind than the brain. Whatever else it is, the brain is also a complex algorithm. But is the brain fully described thereby, or does that omit what makes us distinct, namely, consciousness? Consciousness is the qualitative experience of being somebody or something, its “what-it-is-like-to-be-that”-ness, as one might say. If there is nothing more to the mind than the brain, then algorithms in the era of Big Data will outdo us soon at almost everything we do.¹⁵

A diferença entre a superinteligência e os humanos, na percepção de Risse, se pauta na diferenciação de cérebro e mente, já que o cérebro seria um algoritmo complexo, capaz de aprender e se adaptar, como a superinteligência, mas a mente somente na humanidade é possível. Para ele, a consciência pode ser o que dividirá humanos de robôs¹⁶.

Outra análise filosófica se configura no conceito de singularidade, que seria o momento em que as máquinas superariam os humanos em inteligência¹⁷.

¹⁴STEIBEL, Fabro; VICENTE, Victor Freitas; JESUS, Diego Santos Vieira. Possibilidades e Potenciais da Utilização da Inteligência Artificial. p. 57.

¹⁵ RISSE, Mathias. Human Rights and Artificial Intelligence: An Urgently Needed Agenda. p.3.

¹⁶ RISSE, Mathias. Human Rights and Artificial Intelligence: An Urgently Needed Agenda. p.4.

¹⁷ CHALMERS, David J. **The Singularity: A Philosophical Analysis**. Disponível em: <http://consc.net/papers/singularity.pdf>. Acesso em: 20 jul. 2021.

Since then humans have succeeded in creating something smarter than themselves, this new type of brain may well produce something smarter than itself, and on it goes, possibly at great speed. There will be limits to how long this can continue. But since computational powers have increased rapidly over the decades, the limits to what a superintelligence can do are beyond what we can fathom now. Singularity and superintelligence greatly exercise some participants in the AI debate whereas others dismiss them as irrelevant compared to more pressing concerns.

Tal conceito só foi possível de ser imaginado dado a evolução computacional, mas pode demorar anos ou até décadas. No entanto, se é possível imaginar, como fizemos há anos sobre os robôs baseados em softwares que vemos hoje, talvez seja só uma questão de tempo, não de recursos. Portanto, entende Chalmers, que não pode ser ignorada.¹⁸

Para além do potencial da inteligência artificial, com superinteligência ou não, além de vislumbrar o que ela é capaz ou não de fazer, é preciso entender como ela é internalizada, percebida e sentida pela vida humana.

Por isso, preocupa-se mais qual é sociedade que vamos construir sobre os ombros dos robôs? ¹⁹ É preciso imaginar as pontas de um cabo de guerra entre a lei, a sociedade, o mercado e a tecnologia²⁰.

2. Arcabouço jurídico internacional sobre Inteligência Artificial

Essa preocupação de alinhar mercado, sociedade e tecnologia é global, e reflete todos os relatórios sobre inteligência artificial.

Em 2019, a Harvard Business School²¹ lançou uma organização de todos os guidelines publicados nos últimos 5 anos, que está dividido em 8 temas chave de princípios a serem aplicados em inteligência artificial, compiladas a partir da análise de 36 documentos produzidos pelo setor público, privado e terceiro setor ao redor do mundo.

¹⁸ CHALMERS, David J. **The Singularity**: A Philosophical Analysis.

¹⁹ SOUZA, Carlos Affonso Pereira; OLIVEIRA, Jordan Vinicius. Sobre os ombros de robôs? A inteligência artificial entre fascínios e desilusões. p. 79.

²⁰ SOUZA, Carlos Affonso Pereira; OLIVEIRA, Jordan Vinicius. Sobre os ombros de robôs? A inteligência artificial entre fascínios e desilusões. p. 77.

²¹ FJELD, Jessica; ACHTEN, Nele; HILLIGOSS, Hannah; NAGY, Adam Nagy; SRIKUMAR, Madhulika. **Principled Artificial Intelligence**: Mapping Consensus in Ethical and Rights-based Approaches to Principles for AI. Berkman Klein Center for Internet & Society, 2020.

Os oito temas chaves - *Privacy; Accountability; Safe and Security; Transparency and Explainability; Fairness and Non-Discrimination; Human Control of Technology; Professional Responsibility; Promotion of Human Values*²² - funcionam como categorias nos quais podem ser encontrados vários outras palavras ou conjunto de princípios, valores e fundamentos agrupados sob o termo que consideraram mais amplo.

A *Privacy* contempla os temas *Privacy, Control Over Use of Data, Consent, Privacy by Design, Recommendation for Data Protection Laws, Ability to Restrict Processing, Right to Rectification, Right to Erasure*. Considerando no Brasil a Lei Geral de Proteção de Dados²³ promulgada em 2018, com um processo de vigência bastante confuso. A corrida por dados, por vezes, pode desconsiderar elementos de privacidade e proteção, o que fez com a Europa e o Estado da Califórnia, nos EUA, tivessem legislação específica (GDPR²⁴ e CCPA²⁵).

Já *Accountability* contempla, além da própria palavra da categoria, *Recommendation for New Regulations, Impact Assessment, Evaluation and Auditing Requirement, Verifiability and Replicability, Liability and Legal Responsibility, Ability to Appeal, Environmental Responsibility, Creation of a Monitoring Body, Remedy for Automated Decision*. De forma pragmática, este princípio aparece em exemplos como do carro autônomo, quando tiver que fazer alguma escolha em âmbito moral e/ou ético.

No terceiro tema chave, *Safety and Security*, podem ser encontradas as palavras e expressões *Safety and Reliability, Predictability, Security by Design*. Importante lembrar que qualquer sistema de inteligência artificial vai passar por um processo de calibração, de tentativa e erro. Este processo é normal. Por mais

²²FJELD, Jessica; ACHTEN, Nele; HILLIGOSS, Hannah; NAGY, Adam Nagy; SRIKUMAR, Madhulika. **Principled Artificial Intelligence: Mapping Consensus in Ethical and Rights-based Approaches to Principles for AI**.

²³ BRASIL. **Lei n. 13.709, de 14 de agosto de 2018**. Lei Geral de Proteção de Dados Pessoais (LGPD). Disponível em: http://www.planalto.gov.br/ccivil_03/_ato2015-2018/2018/lei/l13709.htm. Acesso em: 11 jun. 2021.

²⁴ EUROPEAN UNION. **General Data Protection Regulation**. Disponível em: <https://gdpr-info.eu/>. Acesso em: 11 jul. 2021.

²⁵ UNITED STATES OF AMERICA. **California Consumer Privacy Act**. Disponível em: <https://oag.ca.gov/privacy/ccpa>. Acesso em: 11 jul. 2021.

segurança que seja aplicada no processo de desenvolvimento, no começo (e esse tempo é relativo, dependendo de cada aplicação) é preciso um grande controle ou prevenção de dados relativo ao erro.

Para *Transparency and Explainability*, integram *Explainability*, *Transparency*, *Open-Source Data and Algorithms*, *Notification when Interacting with an AI*, *Notification when AI Makes a Decision about an Individual*, *Regular Reporting Requirement*, *Right to Information*, *Open Procurement (for Government)*.

No quinto, *Fairness and Non-discrimination*, as expressões contidas são *Non-discrimination and the Prevention of Bias*, *Fairness*, *Inclusiveness in Design*, *Inclusiveness in Impact*, *Representative and High-Quality Data* e *Equality*. Considerando a questão técnica já levantada, da necessidade de um grande volume de dados para desenvolvimento de inteligência artificial, os tipos de dados e as informações que são utilizadas, é possível inferências discriminatórias. Ou seja, esse princípio tem que ser utilizado não só na criação, mas também no volume de dados.

O sexto, *Human Control of Technology*, as expressões são *Human Review of Automated Decision* e *Ability to Opt out of Automated Decision*. Interessante essa possibilidade de “abrir” a caixa preta da inteligência artificial, no entanto, pode ser uma forma de restrição das técnicas a serem utilizadas.

Em *Professional Responsibility* é possível encontrar *Multistakeholder Collaboration*, *Responsible Design*, *Consideration of Long-Term Effects*, *Accuracy* e *Scientific Integrity*. Por fim, mas não menos importante, *Promotion of Human Values*, que contém *Leveraged to Benefit Society*, *Human Values and Human Flourishing*, *Access to Technology*.

Embora este trabalho seja fundamental para encontrar pontos de convergência, ainda não é possível afirmar que seja unanimidade e muito menos que esteja pacificado. Considerando que este cenário vive uma era efervescente, cabe dizer que, conforme a participação da sociedade e de mais

instituições acontece, será um ambiente ainda mais diversos²⁶. A inteligência artificial vive, ainda, uma era de divergência produtiva.

Para o Direito, litigar sobre as consequências de uma tecnologia ainda nascente em uso, considerando a sua não consolidação, significa lidar com este ambiente em ebulição, e vendo nascer os primeiros julgados, com mais dúvidas do que certezas²⁷. Por isso, esse conjunto de documentos com guias, diretrizes, políticas se apresenta tão fundamental: talvez estejamos diante da discussão basilar sobre as normas internacionais de inteligência artificial sob as quais serão construídas as profissões do futuro, a economia vindoura e a sociedade em que as máquinas têm um papel fundamental.

Nesse sentido, poderão as máquinas nos matar? Relembrando que o botão de liga/desliga é humano, a pergunta que resta a fazer é: como lidar se esse botão passar a ser das máquinas? E se elas forem conscientes, como na superinteligência, desligá-las poderia ser também considerado um crime?

They would have to be designed so they respect human rights even though they would be smart and powerful enough to violate them. At the same time they would have to be endowed with proper protection themselves. It is not impossible that, eventually, the Universal Declaration of Human Rights would have to apply to some of them.²⁸

Risse entende que, eventualmente, a inteligência artificial precisará ser desenvolvida com o princípio asimoviano²⁹ de que não poderá fazer mal a um humano e nem a si mesmo. Eventualmente, até a Declaração de Direitos Humanos precisará ser aplicada à própria inteligência artificial. No entanto, ele mesmo diz que o movimento defensor dos direitos humanos não está preparado para esta discussão especialmente porque, possivelmente, não será capaz de

²⁶ FJELD, Jessica; ACHTEN, Nele; HILLIGOSS, Hannah; NAGY, Adam Nagy; SRIKUMAR, Madhulika. **Principled Artificial Intelligence: Mapping Consensus in Ethical and Rights-based Approaches to Principles for AI.** p.39.

²⁷ FJELD, Jessica; ACHTEN, Nele; HILLIGOSS, Hannah; NAGY, Adam Nagy; SRIKUMAR, Madhulika. **Principled Artificial Intelligence: Mapping Consensus in Ethical and Rights-based Approaches to Principles for AI.** p.39.

²⁸ RISSE, Mathias. Human Rights and Artificial Intelligence: An Urgently Needed Agenda. p.8

²⁹ As 3 leis da robótica previstas por Isac Asimov aludidas nesse texto são:

1) A robot may not injure a human being or, through inaction, allow a human being to come to harm. 2) A robot must obey the orders given it by human beings except where such orders would conflict with the First Law. 3) A robot must protect its own existence as long as such protection does not conflict with the First or Second Laws.

lidar com o resultado³⁰. Para ele, *machines may have to be integrated into human social and political lives*.³¹ Já imaginou que, nessa lógica, poderíamos ter como deputados, senadores, prefeitos e até mesmo presidente da república uma inteligência artificial?

Considerando as revoluções anteriores, a propriedade da terra – na revolução agrícola – e a propriedade das fábricas – na revolução industrial – eram fatores preponderantes (para não dizer determinantes) do poder econômico. Agora, o ativo correspondente, na economia digital, são os dados. Para Risse, *Unequal ownership of data will have detrimental consequences for many people in society as well* ³².

IBM (USA)³³, IA Latam (Chile)³⁴, Telia Company (Suécia)³⁵, Telefónica (Espanha)³⁶, Google (EUA)³⁷, Microsoft (USA)³⁸, ITI (EUA)³⁹, Tencent Institute (China)⁴⁰⁴¹, são empresas privadas, que acabarão atuando, por meio de softwares de inteligência artificial, como entes públicos também, e são as

³⁰ RISSE, Mathias. **Human Rights, Artificial Intelligence and Heideggerian Technoskepticism: The Long (Worrisome?) View**. HKS Faculty Research Working Paper Series RWP19-010, fev. 2019.

³¹ RISSE, Mathias. **Human Rights, Artificial Intelligence and Heideggerian Technoskepticism: The Long (Worrisome?) View**.

³² RISSE, Mathias. **Human Rights, Artificial Intelligence and Heideggerian Technoskepticism: The Long (Worrisome?) View**. p.12.

³³ IBM. **IBM Everyday Ethics for Artificial Intelligence**. 2019. Disponível em: <https://www.ibm.com/watson/assets/duo/pdf/everydayethics.pdf>. Acesso em: 11 jul. 2021.

³⁴ IA LATAM. **Declaración de Ética IA-LATAM para el diseño, Desarrollo y uso de la inteligencia artificial**. 2019. Disponível em: <http://ia-latam.com/etica-ia-latam/>. Acesso em: 11 jul. 2021.

³⁵ TELIA COMPANY. **Guiding Principles on Trusted AI Ethics**. 2019. Disponível em: <https://www.teliacompany.com/globalassets/telia-company/documents/about-telia-company/public-policy/2018/guiding-principles-on-trusted-ai-ethics.pdf>. Acesso em: 11 jul. 2021.

³⁶ TELEFÓNICA. **AI Principles of Telefónica**. 2018. Disponível em: <https://www.telefonica.com/en/web/responsible-business/our-commitments/ai-principles>. Acesso em: 11 jul. 2021.

³⁷ PICHAI, Sundar. AI at Google: our principles. **Google**. Jun. 7 2018. Disponível em: <https://www.blog.google/technology/ai/ai-principles/>. Acesso em: 11 jul. 2021.

³⁸ MICROSOFT. **AI Principles**. 2018. Disponível em: <https://www.microsoft.com/en-us/ai/responsible-ai?activetab=pivot1:primaryr6>. Acesso em: 11 jul. 2021.

³⁹ INFORMATION TECHNOLOGY INDUSTRY COUNCIL. **AI Policy Principles**. 2017. Disponível em: <https://www.itic.org/resources/AI-Policy-Principles-FullReport2.pdf>. Acesso em: 11 jul. 2020.

⁴⁰ TENCENT INSTITUTE. **Six Principles of AI**. 2017. Disponível em: <http://www.kejilie.com/yiyiou/article/ZRZF2.html>. Acesso em: 11 jul. 2021.

⁴¹ FJELD, Jessica; ACHTEN, Nele; HILLIGOSS, Hannah; NAGY, Adam Nagy; SRIKUMAR, Madhulika. **Principled Artificial Intelligence: Mapping Consensus in Ethical and Rights-based Approaches to Principles for AI**. p.39.

empresas que já lançaram algum documento contendo princípios para aplicação de inteligência artificial.

Nesse sentido, também foi assinado um acordo internacional (Arranjo para o Reconhecimento de Critério Comum)⁴² sobre cyber-segurança conduzido pelas empresas privadas de mais de 30 países que tem por objetivo estabelecer bases essenciais para segurança da informação.

Com relação aos governos, já encontramos documentos assinados pela Estados Unidos da América, França⁴³, União Europeia⁴⁴, Reino Unido⁴⁶, Índia⁴⁷, México⁴⁸, Alemanha⁴⁹, Emirados Árabes⁵⁰, Singapura⁵¹, Japão⁵², Bélgica

⁴² CCRA. **Common Criteria Recognition Arrangement**. Disponível em: <https://commoncriteriaportal.org/index.cfm?>. Acesso em: 20 ago. 2021.

⁴³ MISSION ASSIGNED BY THE FRENCH PRIME MINISTER. **For a Meaningful Artificial Intelligence: Toward a French and European Strategy**. 2018. Disponível em: https://www.aiforhumanity.fr/pdfs/MissionVillani_Report_ENG-VF.pdf. Acesso em: 11 jul. 2021.

⁴⁴ EUROPEAN COMMISSION'S HIGH-LEVEL EXPERT GROUP ON ARTIFICIAL INTELLIGENCE. **Ethics Guidelines for Trustworthy AI**. 2018. Disponível em: <https://ec.europa.eu/digital-single-market/en/news/ethics-guidelines-trustworthy-ai>. Acesso em: 11 jul. 2021.

⁴⁵ EUROPEAN COMMISSION. **Artificial Intelligence for Europe**: Communication from the Commission to the European Parliament, the European Council, the Council, the European Economic and Social Committee, and the Committee of the Regions. 2018. Disponível em: <https://ec.europa.eu/digital-single-market/en/news/communication-artificialintelligence-europe>. Acesso em: 11 jul. 2021.

⁴⁶ UK HOUSE OF LORDS, SELECT COMMITTEE ON ARTIFICIAL INTELLIGENCE. **AI in the UK: Ready, Willing and Able?** 2018. Report of Session 2017-19. Disponível em: <https://publications.parliament.uk/pa/ld201719/ldselect/ldai/100/100.pdf>. Acesso em: 11 jul. 2021.

⁴⁷ NITI AAYOG. **National Strategy for Artificial Intelligence: #AI for All** (Discussion Paper). 2018. Disponível em: https://www.niti.gov.in/writereaddata/files/document_publication/NationalStrategy-for-AI-Discussion-Paper.pdf. Acesso em: 11 jul. 2021.

⁴⁸ BRITISH EMBASSY IN MEXICO CITY. **Hacia una estrategia de IA em México: aprovechando la revolución de la IA**. 2018. Disponível em: https://docs.wixstatic.com/ugd/7be025_ba24a518a53a4275af4d7ff63b4cf594.pdf. Acesso em: 11 jun. 2021.

⁴⁹ GERMAN FEDERAL MINISTRY OF EDUCATION AND RESEARCH, THE FEDERAL MINISTRY FOR ECONOMIC AFFAIRS AND ENERGY, AND THE FEDERAL MINISTRY OF LABOUR AND SOCIAL AFFAIRS. **Cabinet decides to updates the federal government's AI strategy**. 2018. Disponível em: <https://www.ki-strategie-deutschland.de/home.html>. Acesso em: 11 jul. 2021.

⁵⁰ SMART DUBAI. **Artificial Intelligence Principles and Ethics**. 2019. Disponível em: <https://smartdubai.ae/initiatives/ai-principles-ethics>. Acesso em: 11 jul. 2021.

⁵¹ MONETARY AUTHORITY OF SINGAPORE. **Principles to Promote Fairness, Ethics, Accountability and Transparency (FEAT) in the Use of Artificial Intelligence and Data Analytics in Singapore's Financial Sector**. 2019. Disponível em: <http://www.mas.gov.sg/-/media/MAS/News%20and%20Publications/Monographs%20and%20Information%20Papers/FEAT%20Principles%20Final.pdf>. Acesso em: 11 jul. 2021.

⁵² JAPANESE CABINET OFFICE, COUNCIL FOR SCIENCE, TECHNOLOGY AND INNOVATION. **Social Principles of Human-Centric Artificial Intelligence**. 2019. Disponível em: <https://www8.cao.go.jp/cstp/english/humancentricai.pdf>. Acesso em: 11 jul. 2021.

e China^{5354, 55}

Fjeld et al ainda criou uma categoria chamada “Multistakeholder” para contemplar outros entes que se pronunciaram sobre o assunto, assim como “Civil Society”. No primeiro grupo, temos AI Industry Alliance (China)⁵⁶, Beijing Academy of AI (China)⁵⁷, New York Times (EUA)⁵⁸, IEEE (EUA)⁵⁹, University of Montreal (Canada) Future of Life Institute (EUA)⁶⁰, Partnership on AI (EUA)⁶¹. No segundo grupo, temos UNI Global Union (Suiça)⁶², Amnesty International⁶³ | Access Now⁶⁴ (Canada), Access Now (EUA), T20: Think20 (Argentina)⁶⁵, The

⁵³ GOVERNANCE Principles for a New Generation of Artificial Intelligence: develop responsible artificial intelligence. **ChinaDaily.com.cn**. 17 jun. 2019. Disponível em: <http://www.chinadaily.com.cn/a/201906/17/WS5d07486ba3103dbf14328ab7.html>. Acesso em: 11 jul. 2021.

⁵⁴ STANDARD ADMINISTRATION OF CHINA. **Translation:** excerpts from China's 'White Paper on Artificial Intelligence Standardization'. Jan. 2018. Disponível em: <https://www.newamerica.org/cybersecurity-initiative/digichina/blog/translation-excerpts-chinas-white-paper-artificial-intelligence-standardization/>. Acesso em: 11 jul. 2021.

⁵⁵ FJELD, Jessica; ACHTEN, Nele; HILLIGOSS, Hannah; NAGY, Adam Nagy; SRIKUMAR, Madhulika. **Principled Artificial Intelligence: Mapping Consensus in Ethical and Rights-based Approaches to Principles for AI**. p.39.

⁵⁶ ARTIFICIAL INTELLIGENCE INDUSTRY ALLIANCE. **Artificial Intelligence Industry Code of Conduct (Consultation Version)**. 2019. Disponível em: <https://www.secrss.com/articles/11099>. Acesso em: 11 jun. 2021.

⁵⁷ BEIJING ACADEMY OF ARTIFICIAL INTELLIGENCE. **Beijing AI Principles**. 2019. Disponível em: <https://www.baai.ac.cn/blog/beijing-ai-principles?categoryId=394>. Acesso em: 11 jun. 2021.

⁵⁸ NEW YORK TIMES. **Seeking Ground Rules for AI**. March 2019. Disponível em: <https://www.nytimes.com/2019/03/01/business/ethical-ai-recommendations.html>. Acesso em: 11 jul. 2021.

⁵⁹ IEEE GLOBAL INITIATIVE ON ETHICS OF AUTONOMOUS AND INTELLIGENT SYSTEMS. **Ethically Aligned Design: A Vision for Prioritizing Human Well-Being with Autonomous and Intelligent Systems**. 2019. Disponível em: https://standards.ieee.org/content/dam/ieee-standards/standards/web/documents/other/ead1e.pdf?utm_medium=undefined&utm_source=undefined&utm_campaign=undefined&utm_content=undefined&utm_term=undefined. Acesso em: 11 jul. 2021.

⁶⁰ FUTURE OF LIFE INSTITUTE. **Asilomar AI Principles**. 2017. Disponível em: <https://futureoflife.org/aiprinciples/?cn-reloaded=1>. Acesso em: 20 jul. 2021.

⁶¹ PARTNERSHIP ON AI. **Tenets**. 2016. Disponível em: <https://www.partnershiponai.org/tenets/>. Acesso em: 11 jul. 2021.

⁶² UNI GLOBAL UNION. **Top 10 Principles for Ethical Artificial Intelligence**. 2017. Disponível em: http://www.thefutureworldofwork.org/media/35420/uni_ethical_ai.pdf. Acesso em: 11 jul. 2021.

⁶³ AMNESTY INTERNATIONAL, ACCESS NOW. **Toronto Declaration: Protecting the Right to Equality and Non-Discrimination in Machine Learning Systems**. 2018. Disponível em: https://www.accessnow.org/cms/assets/uploads/2018/08/The-Toronto-Declaration_ENG_08-2018.pdf. Acesso em: 11 jun. 2021.

⁶⁴ ACCESS NOW. **Human Rights in the Age of Artificial Intelligence**. 2018. Disponível em: <https://www.accessnow.org/cms/assets/uploads/2018/11/AI-and-Human-Rights.pdf>. Acesso em: 11 jun. 2021.

⁶⁵ THINK 20. **Future of Work and Education for the Digital Age**. Disponível em: https://www.g20-insights.org/wp-content/uploads/2019/04/T20-Recommendations-Report_TF7-The-Future-of-Work-and-Education-in-the-Digital-Age.pdf. Acesso em: 11 jul. 2021.

Public Voice Coalition (Bélgica)^{66,67} Por fim, a Organização CDE⁶⁸ e o G20⁶⁹, Grupo das Nações Unidas⁷⁰, também se pronunciaram a respeito.

3. Impacto da Inteligência Artificial no Direito

O Direito é impactado pela tecnologia de duas grandes formas: a primeira, para garantir segurança jurídica das ações realizadas com tecnologia; já a segunda ponta diz respeito a prática jurídica, como será transformada.

Em 2020, a Columbia Law Review lançou um dossiê temático⁷¹ que reflete este cabo de guerra já citado entre lei, sociedade, mercado e tecnologia, dividindo os desafios entre os ramos público e privado. Os aspectos concernentes ao ramo privado não serão escopo deste artigo.

No ramo público, discute-se tanto a possibilidade de leis serem feitas quanto julgadas por inteligência artificial e, neste dossiê, foi cunhado o termo “Cyborg Justice”, que significa a união dos esforços homem-máquina no processo de decisão. Isso porque entende que seres humanos podem ser caros, limitados e arbitrários ao contrário do que seria barato, rápido e escalável com os algoritmos.⁷² Essas limitações humanas tendem a promover o ambiente como um próximo passo natural na direção *a rule of law, not of men*⁷³

⁶⁶ THE PUBLIC VOICE COALITION. **Universal Guidelines for Artificial Intelligence**. 23 out. 2018. Disponível em: <https://thepublicvoice.org/ai-universal-guidelines/>. Acesso em: 11 jul. 2021.

⁶⁷ FJELD, Jessica; ACHTEN, Nele; HILLIGOSS, Hannah; NAGY, Adam Nagy; SRIKUMAR, Madhulika. **Principled Artificial Intelligence: Mapping Consensus in Ethical and Rights-based Approaches to Principles for AI**. p.39.

⁶⁸ ORGANISATION FOR ECONOMIC CO-OPERATION AND DEVELOPMENT. **Recommendation of the Council on Artificial Intelligence**. 2019. Disponível em: <https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0449>. Acesso em: 11 jul. 2021.

⁶⁹ G20 TRADE MINISTERS AND DIGITAL ECONOMY MINISTERS. **G20 Ministerial Statement on Trade and Digital Economy**. 2019. Disponível em: <https://www.mofa.go.jp/files/000486596.pdf>. Acesso em: 11 jul. 2021.

⁷⁰ COUNCIL OF EUROPE, EUROPEAN COMMISSION FOR THE EFFICIENCY OF JUSTICE. **European Ethical Charter on the Use of Artificial Intelligence in Judicial Systems and Their Environment**. 2018. Disponível em: <https://rm.coe.int/ethical-charter-en-for-publication-4-december2018/16808f699c>. Acesso em: 11 jul. 2021.

⁷¹ CLR. **Columbia Law Review** – Artificial Intelligence Archives. Disponível em: https://columbialawreview.org/topic/artificial-intelligence/?post_type=content. Acesso em: 20 jul. 2021.

⁷² CROOTOFF, Rebecca. “Cyborg Justice” and the risk of technological-legal lock-in. **Columbia Law Review**, New York, v. 119, n. 7, p. 233-251. Available at <https://columbialawreview.org/content/cyborg-justice-and-the-risk-of-technological-legal-lock-in/>, p. 237.

⁷³ CROOTOFF, Rebecca. “Cyborg Justice” and the risk of technological-legal lock-in. p. 237.

Cyborg justice systems are proliferating, but their structure has not yet stabilized. We now have a bounded opportunity to design systems that realize the benefits and mitigate the issues associated with incorporating AI into the common law adjudicatory process.⁷⁴

Essa oportunidade de ver o sistema *cyborg* nascer pode ser desafiante e recompensador. Em termos de economia comportamental⁷⁵, interessante já iniciar com os incentivos corretos, assim como com os mapas de escolhas estruturados. Por outro lado, quanto maior a complexidade inicial, maior a dificuldade de conseguir colocar em funcionamento.

Para Crootoof, ambos, humanos e robôs, são *black boxes* que podem sair do controle. Porém, se as limitações humanas colocam os robôs em vantagem, por outro lado, o processo de motivação das decisões escritas por humanos pelo menos apresenta uma justificativa. Isso acontece porque pessoas, por mais que devam desempenhar papéis imparciais, nas palavras de Crootoof, *judges live in the real world and internalize social norms, those norms inform and undergird their reasoning, which in turn strengthens those norms*⁷⁶. Além disso, na mesma literatura, argumenta-se que a percepção de legitimidade da sociedade é maior no caso de humanos, em detrimento da inteligência artificial.

No entanto, essa percepção pode ter prazo determinado. O ser humano aprende por experiência, incluindo, nesse aspecto, a repetição. O próprio conceito de jurisprudência, seja no direito positivado ou no *commom law*, se alimentam dessa repetição. A inteligência artificial não é diferente e, também, aprende por repetição e correção de rota. Ao entender um sistema híbrido que apresentará decisões e será calibrado por um humano sempre que houver um erro, ocorre dois fenômenos interessantes de ser estudado. Ao mesmo tempo que essa calibração subordina a inteligência artificial à supremacia humana e garante a presença, ela também levanta a questão do erro. Isso significa dizer que eles acontecerão, mas qual seria sua dimensão? Quem quer ser o cidadão que recorrerá ao tribunal para a pacificação de seus litígios e receberá como resultado um erro robótico? Erros menores, facilmente calibráveis não seriam

⁷⁴ CROOTOOF, Rebecca. "Cyborg Justice" and the risk of technological-legal lock-in. p. 236.

⁷⁵ SUNSTEIN, Cass; THALER, Richard. **Nudge**: Improving decisions about Health, Wealth and Happiness. London: Penguin Books, 2009.

⁷⁶ CROOTOOF, Rebecca. "Cyborg Justice" and the risk of technological-legal lock-in. p. 238.

um problema, porque o humano soberano estaria solucionando-os. No entanto, e quando isso não acontecer, ou vazar algo que deveria estar em segredo de justiça, não é um risco que se pretende correr.⁷⁷

Partindo do entendimento de que quem pode fazer as leis, também pode julgá-las, o governo dinamarquês criou sete princípios⁷⁸ para a digitalização da lei, preparando-o para ser automatizada e buscar um código (no sentido de programação) para os códigos (no sentido legal)⁷⁹.

O primeiro princípio, *Simple and Clear Rules*, como o próprio nome diz, prega que as leis devem ser claras para que qualquer pessoa ou negócio seja capaz de entender, assim como isso facilitaria para as autoridades administrá-las, inclusive, com ganho de escala tecnológico. Para isso, deve ser retirada toda a ambiguidade e mantida a consistência. Isso não significa, necessariamente, textos menores, mas sim, determinando o que é regra e o que é exceção, para não se perder em um mundo de requisitos que somente um técnico seja capaz de entender. Isso talvez signifique revisitar as leis, especialmente para separar a regra da exceção.

Digital Communication, o segundo princípio, diz que a legislação deve garantir uma comunicação de forma inteiramente digital, seja com pessoas físicas ou jurídicas, com exceção apenas para aqueles que não consigam, fazendo com que métodos alternativos precisem existir em caráter de exceção. Para que isso aconteça, o governo dinamarquês, neste guia, entende que deve existir uma autoridade legal imbuída dessa atribuição, com foco no desenvolvimento tecnológico governamental. Para que possa ser abrangente, por tecnologia, a Dinamarca entende múltiplas plataformas, como SMS, sites, comunicadores instantâneos, e-mail, apps e até mesmo sistemas específicos que se fizerem necessários.

⁷⁷ CROOTOF, Rebecca. "Cyborg Justice" and the risk of technological-legal lock-in. p. 245.

⁷⁸ Regras simples e claras; Comunicação digital; Possibilidade de processamento dos casos automaticamente; Consistência entre autoridades; Manuseio seguro dos dados; Uso de infraestrutura pública e Prevenção de fraudes e erros. (tradução do autor).

⁷⁹ DANISH PARLIAMENT. **Guidance on digital-ready legislation**. Disponível em: <https://en.digst.dk/policy-and-strategy/digital-ready-legislation/guidances-and-tools/seven-principles-for-digital-ready-legislation/>. Acesso em: 20 jul. 2020.

Já o terceiro princípio *Possibility of automated case processing* inclui dizer que a legislação deve apoiar, no todo ou em parte, a administração digital da legislação considerando os direitos do ordenamento jurídico, utilizando critérios objetivos para que não cause dúvidas, retomando a situação em que a lei é clara e simples. No entanto, quando forem situações legislativas em que falte tal clareza, ou mesmo que demande discricção ou segredo de justiça, o escopo de automação deve ser restrito.

Para *Consistency across authorities - uniform concepts and reuse of data*, versa em criar base de dados única, ou pelo menos coesa e consistente, evitando refação e retrabalho. Para isso, uma taxonomia de dados única para todo o governo se faz necessário.

Safe and secure data handling quer dizer alto grau de segurança dos dados, em um cenário com proteção de dados para pessoas naturais e jurídicas. Nesse sentido, se realmente houver benefício para os cidadãos, é possível que eles mesmos queiram divulgar ou fazer parte dessa base, que deve conter todos os esforços técnicos de que não sejam utilizados em prejuízo para os titulares dos dados.

O penúltimo princípio, *Use of public infrastructure*, prega que o ideal é utilizar, sempre que possível, a infraestrutura existente, de preferência, pública, ou em parceria com a privada. Eventualmente, até aquelas de bases *open source*.

Por fim, *Prevention of fraud and errors*, não poderia faltar, já que é necessário prevenir e controlar os dados por meio de ferramentas de tecnologia da informação, monitorando os riscos, considerando a legislação vigente sobre privacidade e proteção de dados, já citada anteriormente.

Em tese, se seguidos estes princípios, um sistema híbrido, ou *cyborg*, estaria mais próximo de existir. Há que se considerar que os juízes humanos, em razão de sua sensibilidade para visualizar a evolução dos comportamentos sociais e seus ordenamentos jurídicos, alinhando-os no tempo, de forma equilibrada, eles não teriam a mesma velocidade e produtividade. Assim, de acordo com Crootoof, talvez esta seja o grande valor dos humanos *to engage in*

*the value balancing and norm incorporation necessary to maintain an evolving and legitimate common law*⁸⁰.

Considerações finais

Estamos diante de um cenário de revolução pautada pelo digital. Silenciosa, talvez, pois tem acontecido ao longo dos últimos vinte anos, pelo menos, como previu Susskind. No entanto, irrefreável, já afirmou Schwab.

Não devemos mais discutir se, no futuro, existirão sistemas de inteligência artificial, mas sim, como eles devem ser integrados a vida política, social e econômica. Terão eles consciência? Terão direitos específicos? Seremos considerados iguais ou haverá alguma supremacia, seja ela humana ou máquina?

O Direito não passará incólume, sequer deve quedar inerte. Faz parte da responsabilidade jurídica se preparar para prover segurança, no ramo privado e público. Isso significa entender o funcionamento, se abrir para o diálogo e, também, ser revolucionado.

Não há espaço para o jurista que não souber interagir com as máquinas, mais do que isso, que conseguir ser produtivo, eficiente e trabalhar em conjunto. Em uma cyborg justice, em que usaremos inteligência artificial para ajudar a escrever leis (*rules as codes*), simplificando o método, garantindo transparência para a sociedade por meio da diferenciação entre regras e exceções; também utilizaremos a inteligência artificial para julgar as situações.

Se as máquinas serão integradas social e politicamente, como defende Risse, naturalmente seria fundamental incorporá-las também ao judiciário. Não como o poder executivo utilizar máquinas de apoio a tomada de decisão e de gestão, o poder legislativo utilizar para escrever leis mais claras e o poder judiciário ser esquecido. Faz parte do processo. As máquinas estarão no cerne do estado democrático de direito.

Cada jurista pode agora tomar somente uma decisão: qual será seu papel nessa revolução, o de protagonista das máquinas ou devorado por elas?

⁸⁰ CROTOF, Rebecca. "Cyborg Justice" and the risk of technological-legal lock-in. p. 251.

Referências

ACCESS NOW. **Human Rights in the Age of Artificial Intelligence**. 2018. Disponível em: <https://www.accessnow.org/cms/assets/uploads/2018/11/AI-and-Human-Rights.pdf>. Acesso em: 11 jun. 2021.

AMNESTY INTERNATIONAL, ACCESS NOW. **Toronto Declaration: Protecting the Right to Equality and Non-Discrimination in Machine Learning Systems**. 2018. Disponível em: https://www.accessnow.org/cms/assets/uploads/2018/08/The-Toronto-Declaration_ENG_08-2018.pdf. Acesso em: 11 jun. 2021.

ARTIFICIAL INTELLIGENCE INDUSTRY ALLIANCE. **Artificial Intelligence Industry Code of Conduct (Consultation Version)**. 2019. Disponível em: <https://www.secrss.com/articles/11099>. Acesso em: 11 jun. 2021.

BEIJING ACADEMY OF ARTIFICIAL INTELLIGENCE. **Beijing AI Principles**. 2019. Disponível em: <https://www.baai.ac.cn/blog/beijing-ai-principles?categoryId=394>. Acesso em: 11 jun. 2021.

BRASIL. **Lei n. 13.709, de 14 de agosto de 2018**. Lei Geral de Proteção de Dados Pessoais (LGPD). Disponível em: http://www.planalto.gov.br/ccivil_03/_ato2015-2018/2018/lei/l13709.htm. Acesso em: 11 jun. 2021.

BRITISH EMBASSY IN MEXICO CITY. **Hacia una estrategia de IA em México: aprovechando la revolución de la IA**. 2018. Disponível em: https://docs.wixstatic.com/ugd/7be025_ba24a518a53a4275af4d7ff63b4cf594.pdf. Acesso em: 11 jun. 2021.

CCRA. **Common Criteria Recognition Arrangement**. Disponível em: <https://commoncriteriaportal.org/index.cfm?>. Acesso em: 20 ago. 2021.

CHALMERS, David J. **The Singularity: A Philosophical Analysis**. Disponível em: <http://consc.net/papers/singularity.pdf>. Acesso em: 20 jul. 2021.

CLR. **Columbia Law Review – Artificial Intelligence Archives**. Disponível em: https://columbialawreview.org/topic/artificial-intelligence/?post_type=content. Acesso em: 20 jul. 2021.

COUNCIL OF EUROPE, EUROPEAN COMMISSION FOR THE EFFICIENCY OF JUSTICE. **European Ethical Charter on the Use of Artificial Intelligence in Judicial Systems and Their Environment**. 2018. Disponível em: <https://rm.coe.int/ethical-charter-en-for-publication-4-december2018/16808f699c>. Acesso em: 11 jul. 2021.

CROOTOF, Rebecca. "Cyborg Justice" and the risk of technological-legal lock-in. **Columbia Law Review**, New York, v. 119, n. 7, p. 233-251, nov. 2019. Disponível em: https://columbialawreview.org/wp-content/uploads/2019/11/Crootof-Cyborg_justice_and_the_risk_of_technological_legal_lock_in.pdf. Acesso em: 11 jul. 2021.

DANISH PARLIAMENT. **Guidance on digital-ready legislation**. Disponível em: <https://en.digst.dk/policy-and-strategy/digital-ready-legislation/guidances-and-tools/seven-principles-for-digital-ready-legislation/>. Acesso em: 20 jul. 2020.

EUROPEAN COMMISSION. **Artificial Intelligence for Europe**: Communication from the Commission to the European Parliament, the European Council, the Council, the European Economic and Social Committee, and the Committee of the Regions. 2018. Disponível em: <https://ec.europa.eu/digital-single-market/en/news/communication-artificialintelligence-europe>. Acesso em: 11 jul. 2021.

EUROPEAN COMMISSION'S HIGH-LEVEL EXPERT GROUP ON ARTIFICIAL INTELLIGENCE. **Ethics Guidelines for Trustworthy AI**. 2018. Disponível em: <https://ec.europa.eu/digital-single-market/en/news/ethics-guidelines-trustworthy-ai>. Acesso em: 11 jul. 2021.

EUROPEAN UNION. **General Data Protection Regulation**. Disponível em: <https://gdpr-info.eu/>. Acesso em: 11 jul. 2021.

FJELD, Jessica; ACHTEN, Nele; HILLIGOSS, Hannah; NAGY, Adam Nagy; SRIKUMAR, Madhulika. **Principled Artificial Intelligence**: Mapping Consensus in Ethical and Rights-based Approaches to Principles for AI. Berkman Klein Center for Internet & Society, 2020.

FUTURE OF LIFE INSTITUTE. **Asilomar AI Principles**. 2017. Disponível em: <https://futureoflife.org/aiprinciples/?cn-reloaded=1>. Acesso em: 20 jul. 2021.

G20 TRADE MINISTERS AND DIGITAL ECONOMY MINISTERS. **G20 Ministerial Statement on Trade and Digital Economy**. 2019. Disponível em: <https://www.mofa.go.jp/files/000486596.pdf>. Acesso em: 11 jul. 2021.

GERMAN FEDERAL MINISTRY OF EDUCATION AND RESEARCH, THE FEDERAL MINISTRY FOR ECONOMIC AFFAIRS AND ENERGY, AND THE FEDERAL MINISTRY OF LABOUR AND SOCIAL AFFAIRS. **Cabinet decides to updates the federal government's AI strategy**. 2018. Disponível em: <https://www.ki-strategie-deutschland.de/home.html>. Acesso em: 11 jul. 2021.

GOVERNANCE Principles for a New Generation of Artificial Intelligence: develop responsible artificial intelligence. **ChinaDaily.com.cn**. 17 jun. 2019. Disponível em: <http://www.chinadaily.com.cn/a/201906/17/WS5d07486ba3103dbf14328ab7.html>. Acesso em: 11 jul. 2021.

IA LATAM. **Declaración de Ética IA-LATAM para el diseño, Desarrollo y uso de la inteligencia artificial**. 2019. Disponível em: <http://ia-latam.com/etica-ia-latam/>. Acesso em: 11 jul. 2021.

IBM. **IBM Everyday Ethics for Artificial Intelligence**. 2019. Disponível em: <https://www.ibm.com/watson/assets/duo/pdf/everydayethics.pdf>. Acesso em: 11 jul. 2021.

IEEE GLOBAL INITIATIVE ON ETHICS OF AUTONOMOUS AND INTELLIGENT SYSTEMS. **Ethically Aligned Design: A Vision for Prioritizing Human Well-Being with Autonomous and Intelligent Systems**. 2019. Disponível em: https://standards.ieee.org/content/dam/ieee-standards/standards/web/documents/other/ead1e.pdf?utm_medium=undefined&utm_source=undefined&utm_campaign=undefined&utm_content=undefined&utm_term=undefined. Acesso em: 11 jul. 2021.

INFORMATION TECHNOLOGY INDUSTRY COUNCIL. **AI Policy Principles**. 2017. Disponível em: <https://www.itic.org/resources/AI-Policy-Principles-FullReport2.pdf>. Acesso em: 11 jul. 2020.

JAPANESE CABINET OFFICE, COUNCIL FOR SCIENCE, TECHNOLOGY AND INNOVATION. **Social Principles of Human-Centric Artificial Intelligence**. 2019. Disponível em: <https://www8.cao.go.jp/cstp/english/humancentricai.pdf>. Acesso em: 11 jul. 2021.

MICROSOFT. **AI Principles**. 2018. Disponível em: <https://www.microsoft.com/en-us/ai/responsible-ai?activetab=pivot1:primaryr6>. Acesso em: 11 jul. 2021.

MISSION ASSIGNED BY THE FRENCH PRIME MINISTER. **For a Meaningful Artificial Intelligence: Toward a French and European Strategy**. 2018. Disponível em: https://www.aiforhumanity.fr/pdfs/MissionVillani_Report_ENG-VF.pdf. Acesso em: 11 jul. 2021.

MONETARY AUTHORITY OF SINGAPORE. **Principles to Promote Fairness, Ethics, Accountability and Transparency (FEAT) in the Use of Artificial Intelligence and Data Analytics in Singapore's Financial Sector**. 2019. Disponível em: <http://www.mas.gov.sg/~media/MAS/News%20and%20Publications/Monographs%20and%20Information%20Papers/FEAT%20Principles%20Final.pdf>. Acesso em: 11 jul. 2021.

NEW YORK TIMES. **Seeking Ground Rules for AI**. March 2019. Disponível em: <https://www.nytimes.com/2019/03/01/business/ethical-ai-recommendations.html>. Acesso em: 11 jul. 2021.

NITI AAYOG. **National Strategy for Artificial Intelligence: #AI for All** (Discussion Paper). 2018. Disponível em: https://www.niti.gov.in/writereaddata/files/document_publication/NationalStrategy-for-AI-Discussion-Paper.pdf. Acesso em: 11 jul. 2021.

ORGANISATION FOR ECONOMIC CO-OPERATION AND DEVELOPMENT. **Recommendation of the Council on Artificial Intelligence**. 2019. Disponível em: <https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0449>. Acesso em: 11 jul. 2021.

PARTNERSHIP ON AI. **Tenets**. 2016. Disponível em: <https://www.partnershiponai.org/tenets/>. Acesso em: 11 jul. 2021.

PICHAU, Sundar. AI at Google: our principles. **Google**. Jun. 7 2018. Disponível em: <https://www.blog.google/technology/ai/ai-principles/>. Acesso em: 11 jul. 2021.

RISSE, Mathias. Human Rights and Artificial Intelligence: An Urgently Needed Agenda. **Human Rights Quarterly**, n.41, v.1, p.1-16, fev. 2019.

RISSE, Mathias. **Human Rights, Artificial Intelligence and Heideggerian Technoskepticism: The Long (Worrisome?) View**. HKS Faculty Research Working Paper Series RWP19-010, fev. 2019.

RISSE, Mathias; LIVINGSTON, Steven. The Future Impact of Artificial Intelligence on Humans and Human Rights. **Ethics and International Affairs**, n.33, v.2, p.141-158, jun. 2019.

SCHWAB, Klaus. **A quarta revolução industrial**. São Paulo: Edipro, 2018.

SMART DUBAI. **Artificial Intelligence Principles and Ethics**. 2019. Disponível em: <https://smartdubai.ae/initiatives/ai-principles-ethics>. Acesso em: 11 jul. 2021.

SOUZA, Carlos Affonso Pereira; OLIVEIRA, Jordan Vinicius. Sobre os ombros de robôs? A inteligência artificial entre fascínios e desilusões. In: FRAZÃO, Ana; MULHOLLAND, Caitlin. **Inteligência Artificial e Direito**. São Paulo: Thompson Reuters, 2019.

STANDARD ADMINISTRATION OF CHINA. **Translation: excerpts from China's 'White Paper on Artificial Intelligence Standardization'**. Jan. 2018. Disponível em: <https://www.newamerica.org/cybersecurity-initiative/digichina/blog/translation-excerpts-chinas-white-paper-artificial-intelligence-standardization/>. Acesso em: 11 jul. 2021.

STEIBEL, Fabro; VICENTE, Victor Freitas; JESUS, Diego Santos Vieira. Possibilidades e Potenciais da Utilização da Inteligência Artificial. In FRAZÃO, Ana; MULHOLLAND, Caitlin. **Inteligência Artificial e Direito**. São Paulo: Thompson Reuters, 2019.

SUNSTEIN, Cass; THALER, Richard. **Nudge**: Improving decisions about Health, Wealth and Happiness. London: Penguin Books, 2009.

SUSSKIND, Richard. **Expert Systems in Law**. London: Oxford University Press, 1987.

SUSSKIND, Richard. **Online Courts and the Future of Justice**. London: Oxford University Press, 2019.

SUSSKIND, Richard. **The Future of the Professions**. London: Oxford University Press, 2017.

SUSSKIND, Richard. **Tomorrow's Lawyers**. London: Oxford University Press, 2017.

TELEFÓNICA. **AI Principles of Telefónica**. 2018. Disponível em: <https://www.telefonica.com/en/web/responsible-business/our-commitments/ai-principles>. Acesso em: 11 jul. 2021.

TELIA COMPANY. **Guiding Principles on Trusted AI Ethics**. 2019. Disponível em: <https://www.teliacompany.com/globalassets/telia-company/documents/about-telia-company/public-policy/2018/guiding-principles-on-trusted-ai-ethics.pdf>. Acesso em: 11 jul. 2021.

TENCENT INSTITUTE. **Six Principles of AI**. 2017. Disponível em: <http://www.kejilie.com/iyiou/article/ZRZFn2.html>. Acesso em: 11 jul. 2021.

THE PUBLIC VOICE COALITION. **Universal Guidelines for Artificial Intelligence**. 23 out. 2018. Disponível em: <https://thepublicvoice.org/ai-universal-guidelines/>. Acesso em: 11 jul. 2021.

THINK 20. **Future of Work and Education for the Digital Age**. Disponível em: https://www.g20-insights.org/wp-content/uploads/2019/04/T20-Recommendations-Report_TF7-The-Future-of-Work-and-Education-in-the-Digital-Age.pdf. Acesso em: 11 jul. 2021.

UK HOUSE OF LORDS, SELECT COMMITTEE ON ARTIFICIAL INTELLIGENCE. **AI in the UK: Ready, Willing and Able?** 2018. Report of Session 2017-19. Disponível em: <https://publications.parliament.uk/pa/ld201719/ldselect/ldai/100/100.pdf>. Acesso em: 11 jul. 2021.

UNI GLOBAL UNION. **Top 10 Principles for Ethical Artificial Intelligence**. 2017. Disponível em: http://www.thefutureworldofwork.org/media/35420/uni_ethical_ai.pdf. Acesso em: 11 jul. 2021.

UNITED STATES OF AMERICA. **California Consumer Privacy Act**. Disponível em: <https://oag.ca.gov/privacy/ccpa>. Acesso em: 11 jul. 2021.