

# DISCRIMINAÇÃO ALGORÍTMICA, INTELIGÊNCIA ARTIFICIAL, HIPERVIGILÂNCIA DIGITAL E TOMADA DE DECISÃO AUTOMATIZADA

**Haide Maria Hupffer  
Wilson Engelmann  
Gabriel Cemin Petry  
Juliane Altmann Berwig  
(organizadores)**

Este livro é o resultado da pesquisa e das relações interinstitucionais produzidas no âmbito do seguinte projeto de investigação científica:

INTELIGÊNCIA ARTIFICIAL E SOCIEDADE DE ALGORITMOS: REGULAÇÃO, RISCOS DISCRIMINATÓRIOS, GOVERNANÇA E RESPONSABILIDADES

Fundação de Amparo à Pesquisa do Estado do Rio Grande do Sul - FAPERGS

Processo número  
61271.674.22274.26112021,

Edital Fapergs 07/2021  
Programa Pesquisador Gaúcho – PqG



**HAIDE MARIA HUPFFER  
WILSON ENGELMANN  
GABRIEL CEMIN PETRY  
JULIANE ALTMANN BERWIG  
(ORGANIZADORES)**

**DISCRIMINAÇÃO ALGORÍTMICA,  
INTELIGÊNCIA ARTIFICIAL,  
HIPERVIGILÂNCIA DIGITAL  
E TOMADA DE DECISÃO AUTOMATIZADA**

CASA LEIRIA  
SÃO LEOPOLDO/RS  
2024

# DISCRIMINAÇÃO ALGORÍTMICA, INTELIGÊNCIA ARTIFICIAL, HIPERVIGILÂNCIA DIGITAL E TOMADA DE DECISÃO AUTOMATIZADA

Organizadores: Haide Maria Hupffer  
Wilson Engelmann  
Gabriel Cemin Petry  
Juliane Altmann Berwig

DOI: <https://doi.org/10.29327/5448881>

Capa: baseada em imagem criada por IA Image Gen App.

Os textos são de responsabilidade de seus autores.

Qualquer parte desta publicação pode ser reproduzida, desde que citada a fonte.

## Casa Leiria Conselho Editorial

|                                          |            |
|------------------------------------------|------------|
| Ana Carolina Einsfeld Mattos             | (UFRGS)    |
| Ana Patrícia Sá Martins                  | (UEMA)     |
| Antônia Sueli da Silva Gomes Temóteo     | (UERN)     |
| Glícia Marili Azevedo de Medeiros Tinoco | (UFRN)     |
| Haide Maria Hupffer                      | (Feevale)  |
| Isabel Cristina Arendt                   | (Unisinos) |
| Isabel Cristina Michelan de Azevedo      | (UFS)      |
| José Ivo Follmann                        | (Unisinos) |
| Luciana Paulo Gomes                      | (Unisinos) |
| Luiz Felipe Barboza Lacerda              | (UNICAP)   |
| Márcia Cristina Furtado Ecoten           | (Unisinos) |
| Rosângela Fritsch                        | (Unisinos) |
| Tiago Luís Gil                           | (UnB)      |

D611 Discriminação algorítmica, inteligência artificial, hipervigilância digital e tomada de decisão automatizada [recurso eletrônico]. / organização Haide Maria Hupffer...[et.al]. – São Leopoldo: Casa Leiria, 2024.

Disponível em: <<http://www.casaleiriaacervo.com.br/direito/discriminacao/index.html>>

Conteúdo: textos em português, inglês e espanhol.

ISBN 978-85-9509-140-5

1. Direito – Tecnologia. 2. Direitos humanos – Tecnologia. 3. Direito – Inteligência artificial – Impactos legais, éticos e sociais. 4. Direito – Tecnologia – Justiça, transparência e responsabilidade. I. Hupffer, Haide Maria (Org.).

CDU 34:004

## SUMÁRIO

|     |                                                                                                                                                                                                                                                     |
|-----|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 9   | Apresentação<br><i>Haide Maria Hupffer</i><br><i>Wilson Engelmann</i><br><i>Gabriel Cemin Petry</i><br><i>Juliane Altmann Berwig</i>                                                                                                                |
| 11  | Hipervigilância, hacking governamental, genocídio quântico e <i>big data</i><br><i>Paola Cantarini</i>                                                                                                                                              |
| 41  | Discriminación algorítmica: un análisis jurídico de los desafíos y oportunidades en la era digital<br><i>Pablo Rafael Banchio</i>                                                                                                                   |
| 69  | Problemas jurídicos para controlar la discriminación algorítmica<br><i>José Julio Fernández Rodríguez</i>                                                                                                                                           |
| 117 | Inteligência Artificial: um “artefato” discursivo<br><i>Alejandro Knaesel Arrabal</i><br><i>Wilson Engelmann</i>                                                                                                                                    |
| 131 | Do ciberespaço ao capitalismo de vigilância: uma nova configuração de poder arquitetada sobre dados<br><i>Andressa Kerschner</i><br><i>Haide Maria Hupffer</i>                                                                                      |
| 159 | (Des)igualdade, viés e discriminação algorítmica: os direitos humanos entre equidade, imparcialidade e não-discriminação<br><i>André Olivier</i>                                                                                                    |
| 183 | Discriminação algorítmica e LGPD: o papel do princípio da adequação e qualidade nos <i>datasets</i> de treinamento de Sistemas de IA<br><i>Gabriel Cemin Petry</i><br><i>João Sérgio dos Santos Soares Pereira</i><br><i>Karin Regina Rick Rosa</i> |
| 209 | A morte da privacidade na era da hipervigilância digital<br><i>Francisco Soares Campelo Filho</i>                                                                                                                                                   |

- 233 Algoritmos como mecanismos de opressão à população LGBTQIAPN+  
*Ana Paola de Castro e Lins*  
*José Anchieta Oliveira Feitoza*
- 251 Hyper Surveillance In The Criminal Justice System: Balancing Security And Civil Liberties In South Asian Countries (Afghanistan, Bangladesh, Bhutan, India, Maldives, Nepal, Pakistan, and Sri Lanka)  
*Ratna Sisodiya*  
*Rachana Choudhary*  
*Surendra Singh Bhati*
- 273 O devido processo digital na perspectiva processual e constitucional: para além da transparência algorítmica  
*Alana Gabriela Engelmann*  
*Wilson Engelmann*
- 291 A precaução emanada de recomendações internacionais para edificação de estratégias responsáveis no desenvolvimento e implementação de sistemas de inteligência artificial  
*Rafael Cemin Petry*  
*Juliane Altmann Berwig*
- 315 Sesgos en la era digital: desigualdad algorítmica y pueblos indígenas  
*Renato Sebastiani-León Mazza*  
*María Eugenia-Zevallos Loyaga*
- 345 Navegando no metaverso: benefícios, riscos, discriminação algorítmica e a necessidade de regulamentação para proteger o consumidor  
*Tauane da Silva Brito*  
*Haide Maria Hupffer*
- 373 Autoria e inteligência artificial  
*Maiara Ubinski Bauer*  
*André Rafael Weyermüller*
- 403 A plataformação da desinformação e o Projeto de Lei das Fake News  
*Agnes Borges Kalil*  
*Haide Maria Hupffer*  
*Dailor dos Santos*

- 433 Reconhecimento facial e política criminal: entre a segurança pública e a proteção aos direitos fundamentais  
*Luan Tomaz Ferreira*  
*Daniel Kessler de Oliveira*  
*Diogo Machado de Carvalho*
- 467 Inteligência artificial e a prova no processo penal: aproximações teóricas  
*Adilson José Bressan*  
*Reginaldo Pereira*
- 487 Índice remissivo

# **A PRECAUÇÃO EMANADA DE RECOMENDAÇÕES INTERNACIONAIS PARA EDIFICAÇÃO DE ESTRATÉGIAS RESPONSÁVEIS NO DESENVOLVIMENTO E IMPLEMENTAÇÃO DE SISTEMAS DE INTELIGÊNCIA ARTIFICIAL<sup>1</sup>**

*Rafael Cemin Petry<sup>2</sup>*  
*Juliane Altmann Berwig<sup>3</sup>*

## **1. INTRODUÇÃO**

A inteligência artificial (AI) a cada dia está mais presente no cotidiano de todas as pessoas do mundo e, ao mesmo tempo, a cada

- 
- 1 Este trabalho é o resultado parcial da pesquisa realizada no âmbito dos projetos de investigação científica: “Inteligência Artificial e Sociedade de Algoritmos: Regulação, Riscos Discriminatórios, Governança e Responsabilidades” – Edital Fapergs 07/2021 – Programa Pesquisador Gaúcho – PqG.
  - 2 Bacharel em Filosofia pelo Centro Universitário de Grandes Dourados; Integrante do Grupo de Pesquisa CNPq/Feevale Direito e Desenvolvimento; Bolsista IC do Projeto Novas Tecnologias e Sociedade de Risco, bacharelando no curso de Ciências Jurídicas na Universidade Feevale.
  - 3 Doutora em Direito pela Universidade do Vale do Rio dos Sinos com Bolsa pela Coordenação de Aperfeiçoamento de Pessoal de Nível Superior (Capes) pelo Programa de Excelência Acadêmica (Proex). Mestre em Direito pela Universidade do Vale do Rio dos Sinos, especialista em Direito Ambiental Nacional e Internacional pela Universidade Federal do Estado do Rio Grande do Sul e graduada em Direito pela Universidade de Santa Cruz do Sul. Professora no curso de Direito da Universidade FEEVALE e Pesquisadora com projeto financiado pela Fapergs “A INTELIGÊNCIA ARTIFICIAL COMO FERRAMENTA PARA O FRAMEWORK DA GESTÃO DOS RISCOS AMBIENTAIS DAS NANOTECNOLOGIAS”. Orientadora de bolsistas de iniciação científica. Professora de cursos de cursos de Especialização. Participou do Workshop sobre Public Law & Policy na Universidade de Berkeley (EUA), do debate STS Circle at Harvard do Programa em Ciência, Tecnologia e Sociedade na Universidade de Harvard na escola John F. Kennedy, Boston (EUA) e da Conferência Climate Lecture Series na Universidade MIT, Boston (EUA). Fundadora e diretora da Associação Gaúcha dos Advogados de Direito Ambiental Empresarial – AGAAE. Autora de artigos científicos, capítulos de livros e livros. Advogada.

dia mais dela se depende. Diante da dinamicidade assustadora do desenvolvimento das IA muitas dúvidas surgem sobre seus riscos, já que as oportunidades são claramente visíveis e estampadas pelas empresas desenvolvedoras dos sistemas. A questão é o outro lado, onde residem os riscos humanos, às questões éticas, às ameaças aos Direitos Humanos.

A partir deste cenário o presente capítulo tem como objetivo geral: abordar a precaução emanada de recomendações internacionais para edificação de estratégias responsáveis no desenvolvimento e implementação de sistemas de inteligência artificial. Diante disso, o problema de pesquisa revela-se no seguinte questionamento: Quais são as facetas do Princípio da Precaução voltado para as estratégias de proteção humana em sistemas de IA?

Para atender a meta de pesquisa parte-se inicialmente da abordagem sobre o surgimento das IA e seus desdobramentos, na sequência abordar-se-á o princípio da precaução como medida cautelar para os efeitos das tecnologias emergentes e, por fim, os reflexos da precaução com os riscos da ia nas organizações internacionais.

A metodologia da pesquisa adotada é a exploratória, com método de abordagem dedutivo, utilizando-se da pesquisa bibliográfica como procedimento técnico.

## **2. O SURGIMENTO DA INTELIGÊNCIA ARTIFICIAL E SEUS DESDOBRAMENTOS**

Anteriormente a aparição de máquinas capazes de simular a cognição humana através de sistemas computadorizados, pensadores da antiguidade à modernidade já se questionavam acerca da possibilidade de existir similitude entre computar e racionalizar, tendo-se como exemplo o corporeísmo mecanicista hobbesiano (Antiseri; Reale, 2017, p. 390). O filósofo tratava, em suas lições, de alçar uma notória antítese as premissas aristotélicas prescritas no tratado *metafísica* quanto a finalidade da filosofia, desvencilhando-a das puras atividades contemplativas. Para Hobbes, a tarefa e o objeto da filosofia são direcionados aos corpos, os dividindo como i) corpos naturais inanimados, ii) naturais animados, como o homem, e iii) corpos artificializados, sendo um dos representantes mais conhecidos as teorias do Estado (Antiseri; Reale, 2017, p. 393).

Denota-se que preceitos metafísicos, espirituais e desvencilhados do plano físico estão excluídos de sua linha filosófica. A inteligência e os pensamentos emitidos pelo homem, para o empirista, derivam de uma lógica edificada em regras para o correto pensar. Pensamentos seriam fluidos e por esta razão deveriam ser relacionados com sinais sensíveis, pontos de definição e referência, tanto positivos quanto negativos, como “homem”, “não-homem”, “animal”, “não-animal” (Antiseri; Reale, 2017, p. 394). Assim racionalizar seria uma atividade humana ligada a computação ou conexão de nomes, proposições ou definições, caminhando em um sentido embrionário do que viria posteriormente a ser conhecido como cibernética:

Essa concepção do raciocinar, entendido como “compor” e “decompor” e “recompor” e baseado nos semantemas ou sinais linguísticos, e o relativo pano de fundo convencionalista surpreendem pela modernidade e extraordinária audácia, pois contém pressentimentos da contemporânea cibernética (Antiseri; Reale, 2017, p. 395).

Compreensão similar ao estrito sentido entre “pensar” e “computar” pode ser verificado nas teorias do criptoanalista Alan Turing, questionando, em meados de 1936, se máquinas seriam capazes de possuir inteligência. O objetivo acadêmico de Alan Turing vislumbra-se em seus artigos *On Computable Numbers* e *Computing Machinery and Intelligence*, em que empreende esforços para conceituar o que seria uma máquina inteligente, assim como dispõe acerca de maneiras para testá-la (Taulli, 2020, p. 17). Nesta oportunidade, emerge o conhecido Teste de Turing (TT), em que uma máquina, treinada para análises sintáticas e para formalização de respostas simbólicas, confecciona respostas para um entrevistador humano, pairando a tarefa a um avaliador, também humano, de distinguir se o entrevistador dialoga com uma máquina ou outro humano (Taulli, 2020, p. 17).

Conforme Paulo Antônio Caliendo Velloso da Silveira (2021, p. 31), Alan Turing, ao deparar-se com a impossibilidade em distinguir com precisão os conceitos entre inteligência e atividade pensante, o criptoanalista encontra uma forma de testar se uma máquina estaria habilitada a simular o comportamento comunicativo humano, podendo assim confundir um observador neutro. Observando-se o teste edificado por Alan Turing, se verifica a imperiosidade do tra-

tamento de informações como ferramenta fundamental para uma máquina simular o comportamento humano, analisa-se:

O que é genial nesse conceito é que não há necessidade de verificar se a máquina realmente sabe algo, é auto-consciente ou mesmo se está correta. Em vez disso, o teste de Turing indica que uma máquina pode processar grandes quantidades de informações, interpretar a fala e comunicar-se com seres humanos (Taulli, 2020, p. 17).

As teorias do cientista da computação trouxeram provocações e críticas a respeito do paradigma levantado a respeito da possibilidade de máquinas serem habilitadas a simular a cognição humana, incitando produções acadêmicas e científicas acerca da inteligência computacional, marcando-se os anos ulteriores a suas teses (Silveira, 2021, p. 41). Em meados de 1950, por sua vez, observou-se os reflexos da empreitada instaurada por Turing, emergindo no campo da ciência da computação pesquisas com o objetivo de recriar, com regras de precisão e clareza procedimental, a inteligência humana numa máquina, surgindo nomes como Herbert Simon, Marvin Minsky e John McCarthy (Lee, 2019, p. 19). No mesmo sentido, harmoniza com o entendimento acima as premissas do filósofo da informação Pierre Lévy (2022, p. 3) acerca do conceito de Inteligência Artificial:

Desde a década de 1950, o ramo específico da ciência da computação que se preocupa com a modelagem e simulação de inteligência humana é chamado de Inteligência Artificial. A modelagem computacional da inteligência humana é um objetivo científico digno, que teve, e continuará a ter, consideráveis benefícios teóricos e práticos (Lévy, 2022, p. 3).

As premissas do filósofo deslocam-se ao encontro de eventos históricos, como atentou-se na Conferência Acadêmica da Universidade de Dartmouth de 1956, organizada pelo cientista da computação John McCarthy, sendo denominada como “Um estudo da Inteligência Artificial” (Lévy, 2022, p. 21). O estudo objetivou a atribuição artificial de características de aprendizagem em uma máquina, buscando-se um detalhamento de maneira precisa para que a máquina pudesse simular os aspectos de inteligência no uso da linguagem, a habilitando para resolução de problemas anteriormente só solucionáveis pela mão humana (Taulli, 2020, p. 23).

Desta vertente científica e acadêmica emergiram contribuições para sedimentação de métodos de aprendizagem computacional em i) regras, como sistemas simbólicos codificáveis em regras lógicas e ii) redes neurais, buscando-se a edificação de um sistema de aprendizagem similar ao do cérebro humano (Lee, 2019, p. 20). A empreitada científica, ulteriormente com as contribuições de Geoffrey Hinton, tornou possível potencializar o desenvolvimento e processamento das redes neurais para criar o *deeplearning*, método de aprendizagem profunda, capacitando os sistemas de Inteligência Artificial para executar com maior celeridade e precisão tarefas como reconhecimento de fala ou identificação de objetos (Kaufman, 2020, p. 9), despertando-se, consoante entreluz Kai-fu Lee (2019, p. 23), o interesse de pesquisadores, CEOs de tecnologia e futuristas quanto a capacidade da Inteligência Artificial em resolver problemas do mundo real, concomitantemente com a possibilidade de aplicações para reconhecimento de imagens, prever comportamentos, decifrar a fala, tomar decisões e até mesmo efetuar a condução de veículos automobilísticos, impactando campos societários de execução manual, bem como aspectos das tecnologias de informação e comunicação (TICs):

O vapor alterou fundamentalmente a natureza do trabalho manual e a TIC fez o mesmo com certos tipos de trabalho cognitivo. A IA vai modificar ambos. Realizará muitos tipos de tarefas físicas e intelectuais com uma velocidade e potência que ultrapassarão em muito qualquer ser humano, aumentando drasticamente a produtividade em tudo, desde o transporte até a manufatura e a medicina (Lee, 2019, p. 181).

Contemporaneamente, ante a contínua exploração para o desenvolvimento técnico em áreas como Inteligência Artificial e robótica, torna-se possível encontrar, a título de exemplo, diversos tipos de veículos autônomos instalados com sensores de alta tecnologia, como aviões, caminhões, barcos e drones, expandindo-se também as aplicações destes veículos em uma diversidade de processos, como a agricultura de precisão (Schwab, 2016, p. 25). A utilização de drones torna-se um exemplo do refinamento tecnológico para equipamentos capazes de sentir e responder ao ambiente que estão inseridos, capacitando-os para evitar choques e colisões, tendo aplicações em

diversos setores, entre eles, a agricultura para aplicação de adubo e irrigação (Schwab, 2016, p. 24).

A correspondência e interação das soluções em IA com ambientes externos se deve, conforme refere Byung-Chul Han (2021, p. 93) a paulatina transparência das coisas do mundo real, lecionando que esta transparência do mundo ocorre na medida em que é digitalizado e transformado em dados, exemplificando que tal processo não ocorre apenas com o mundo, mas na mesma medida com o próprio comportamento humano, que vem progressivamente ficando transparente, na medida em que algoritmos e a Inteligência Artificial também o tornam calculável e controlável.

Nesta esteira, o filósofo apresenta duas ordens com pertinência a temática, a i) ordem terrena, vinculada a preservação da terra e seus segredos como algo essencialmente indisponível à sede de exploração da mão humana, que tende a tomá-la como recurso ser incessantemente explorado, e a ii) ordem digital, que embasada pelo *habitus digital*, almeja tornar tudo imediatamente disponível, onde dissipam-se os segredos em prol da plena transparência através de sua transformação em dados calculáveis, controláveis e manuseáveis, que consomem todos os pressupostos para uma ordem da terra (Han, 2021, p. 93).

O cuidado com a exploração da técnica, no que tange a impossibilidade de ter-se uma precisa ciência a respeito de seus efeitos e impactos futuros está contida nas linhas filosóficas de Hans Jonas (Oliveira, 2014, p. 128), firmando que os males estão inexoravelmente presentes nos riscos de qualquer ação que orbite sobre a aplicação da técnica, mesmo que seja empreendida com fins preliminarmente benéficos. Para o filósofo, o longo prazo dos efeitos que caminham conjuntamente com a magnitude do poder da tecnologia conduz a um panorama de incerteza sobre as consequências das aplicações da técnica *a posteriori*, importando, com a incerteza dos efeitos, em risco digno de preocupação (Oliveira, 2014, p. 128).

Diante deste cenário de incertezas quanto aos avanços da técnica, Hans Jonas busca sedimentar uma ética como forma de um “poder sobre o poder”, com foco em exercícios prognósticos e profiláticos para vigiar o poder da técnica, ao mesmo passo que se considera as incertezas de seus impactos (Oliveira, 2014, p. 127). Assim, o filósofo apresenta o Princípio da Responsabilidade como proposta

ética para preservação transgeracional da vida, através da i) futurologia comparativa, ii) heurística do temor e iii) imposição de freios voluntários para a exploração exponencial da técnica, com atenção especial ao meio ambiente, observa-se:

O homem deve descobrir que é preciso transpassar o abismo que o separou da natureza. Ele precisa ser reintegrado aquele horizonte de conexões vitais das quais faz parte e com as quais partilha um destino comum. [...] Uma ética com dupla responsabilidade, portanto: proteger a natureza e proteger as gerações futuras. Ou, numa única obrigação moral: Respeitar e cuidar da vida em geral para que também a humana continue sendo possível no futuro (Oliveira, 2014, p. 142).

Assim, como verificou-se, os esforços teóricos, acadêmicos e científicos acerca da possibilidade de uma máquina simular a capacidade cognitiva humana corroboraram para formalização dos sistemas de Inteligência Artificial que atualmente fazem parte do dia a dia de inúmeras pessoas ao redor do mundo, numa pluralidade de aplicações e ações algorítmicas que se relacionam a quase todos os setores da vida em sociedade (Taulli, 2020, p. 13). Entretanto, a inovação tecnológica não deve andar desamparada de critérios de condução ética, visto que a forma com que pode ser administrada detém potencial para provocar impactos incertos não apenas a geração atual, mas a todas as outras que se sucederão.

### **3. O PRINCÍPIO DA PRECAUÇÃO COMO MEDIDA CAUTELAR PARA OS EFEITOS DAS TECNOLOGIAS EMERGENTES**

Em meados de 1959, com Lei de Energia Atômica (Atomgesetz) o Princípio da Precaução passou a ser conhecido, a partir de então, seu reconhecimento mundial vem sendo pulverizado pelo mundo. Ocorre que, todavia, a sua aplicação como medida cautelar e prática sempre foi alvo de discussões (Berwig; Engelmann, 2023). Desde a revolução industrial, percebe-se que o Princípio da Precaução tem sido amplamente utilizado pela Doutrina e legislações internacionais como instrumento de proteção aos danos decorrentes da mecanização, das tecnologias, das inovações. Em cenários de maiores “incertezas” científicas sempre pairam dúvidas sobre os “recortes” e aplicações concretas do princípio, tanto que Sunstein questiona:

Se você levasse o princípio ao máximo, nunca poderia atravessar a rua (afinal, você poderia ser atropelado por um carro); você nunca poderia passar por cima de uma ponte (afinal, ela poderia desabar); você nunca poderia se casar (afinal, isso poderia acabar em um desastre), etc. Se alguém realmente agisse dessa maneira, logo acabaria em uma instituição para doentes mentais (Sunstein, 2005, p.109-110).

Ou seja, o autor retrata a real dificuldade de aplicação da Precaução, dado que os elementos que suportam a sua aplicação são abstratos, culturais, complexos e incertos. Considerando estas complexidades na materialização do princípio em aplicações concretas, é indicado na Declaração Rio 92 que o princípio da precaução tem em sua teleologia o cuidado para com o meio ambiente, devendo ser observado pelos Entes Federativos em harmonização com suas próprias capacidades, consoante seu art. 15 da Declaração: “Quando houver ameaça de danos graves ou irreversíveis, a ausência de certeza científica absoluta não será utilizada como razão para o adiamento de medidas economicamente viáveis para prevenir a degradação ambiental.”.

Veja-se que a composição do Princípio da Precaução perpassa por 03 elementos: i) ameaça de danos graves ou irresistíveis, que são os riscos atinentes as atividades; ii) ausência de certeza científica absoluta; iii) medidas economicamente viáveis para prevenir a degradação ambiental. Ou seja, todos estes elementos serão variáveis, pois dependerão de outros parâmetros, locais, culturais e, portanto, instáveis.

Sunstein argumenta que o Princípio da Precaução é incoerente. Os riscos existem em todos os lados das situações sociais e as medidas de precaução criam os seus próprios perigos. Diversas culturas centram-se em riscos muito diferentes, muitas vezes porque as influências sociais e as pressões dos pares acentuam alguns receios e reduzem outros (Sunstein, 2005, p.109-110). O autor critica duramente o Princípio da Precaução, mas não descarta a importância de se tomar precauções. Uma apreciação dos objetivos salutares do Princípio da Precaução abre caminho para uma reconstrução em três dimensões. O primeiro envolve riscos catastróficos. A segunda envolve danos irreversíveis. A terceira envolve margens de segurança para riscos que, embora não sejam potencialmente catastróficos,

apresentam razões distintas para preocupação. (Sunstein, 2005, p.109-110)

Veja-se, portanto, que é necessário traçar diretrizes para a aplicabilidade da Precaução, as quais de acordo com Berwig e Engelmann (2023), são decorrentes de “uma análise criteriosa dos seguintes pressupostos: i) Risco; ii) Estado da técnica; e iii) Medidas economicamente viáveis”.

O risco perpassa pela **Percepção de Risco** decorrente do conhecimento técnico-científico, mas também é decisiva a participação do senso comum; a **Análise de Risco** com base em um conhecimento técnico-científico especializado; e a **Gestão de Risco**, etapa final de tomada de decisões, com base nos resultados fornecidos pela Análise de Risco (Cezar; Abrantes, 2003). Quanto ao estado da técnica diante a dinamicidade da evolução do conhecimento científico as orientações sobre a Comunicação da Comissão Relativa ao Princípio da Precaução devem considerar a sua “i) Temporalidade pois irá durar enquanto houver incerteza, alterando-se na medida em que novas descobertas do conhecimento científico ocorram; ii) Proporcionalidade, não sendo exigido mais do que a adequação entre o meio utilizado e o fim desejado, avaliando as circunstâncias do caso em concreto. Neste sentido também as iii) Medidas Economicamente Viáveis, como um terceiro elemento da precaução, refere que o custo deve ser ponderado de acordo com a realidade econômica de cada localidade (Berwig; Engelmann, 2023).

Sunstein ataca a Precaução em relação a ideia de que os reguladores devem tomar medidas para proteger contra danos potenciais, mesmo que as cadeias causais sejam incertas e mesmo que não saibamos que os danos são susceptíveis de se concretizarem, ou seja, devam tomar decisões mesmo em um cenário de extrema abstraldade (Sunstein, 2005, p.109-110)

Importante salientar que, apesar de todas as dificuldades de aplicação empírica do Princípio da Precaução, ele permanece sendo o mais importante instrumento de proteção aos riscos e adversidades das novas tecnologias, especialmente da Inteligência Artificial, tema do presente artigo.

Portanto, riscos e contingências devem ser discutidos hoje visando um preparo para o futuro. Também, a partir das novas tecnologias, tal como a Inteligência Artificial, a necessidade de gestão

dos riscos se faz essencial e marca o início de uma nova era que demonstra a necessidade de gerenciamento das tecnologias emergentes a partir da inteligência humana para lidar com problemas de alta complexidade (Berente, 2021, p. 1433-1450).

A Inteligência Artificial (IA) atualmente possui a capacidade de realizar determinadas tarefas com maior velocidade e eficiência em comparação aos seres humanos. Assim, na tomada de decisão, contribuem com sua alta capacidade de processar numerosos dados e informações complexas, de forma mais precisa que a mente humana, podendo para tanto criar cenários, analisar riscos, possibilidades de ocorrência, auxiliando assim em alternativas para solucionar questões de forma estratégica. Ademais, a IA já é usada para detecção e prevenção de desastres naturais ações de combate ao aquecimento global. Portanto, a IA assim como as demais novas tecnologias está atrelada tanto na criação de novos riscos quanto no auxílio para o gerenciamento dos riscos já existentes ou futuros. Mas o Princípio da Precaução, de que forma ele se revela em ações quanto a Inteligência Artificial?

Beck (2018, p. 11), menciona que as mudanças ocorridas no mundo decorrentes da modernidade provocam um “choque fundamental, uma alteração que rompe as constantes antropológicas de nossa existência e de nossas compreensões anteriores de mundo”. Ele menciona assim que o mundo de hoje perpassa por uma metamorfose em que “o que foi impensável ontem é real e possível hoje”. Ou seja, o desenvolvimento da IA era impensável no passado e hoje é presente diariamente na vida de todos.

Percebe-se que diferentemente da aplicação da precaução quanto as nanotecnologias, por exemplo, em que as ações estão relacionadas a gama de resultados científicos de riscos apresentados por produtos, substâncias e mecanismos, na IA as ações de precaução circundam mais em medidas éticas, protetivas aos Direitos Humanos (vida, liberdade, honra, privacidade, emprego...).

Neste sentido, Schwab (2018, p. 188) em relação a IA menciona quais seriam os maiores problemas que se deve considerar:

- Normas éticas: é necessária a criação de princípios e diretrizes relativos aos padrões éticos e expectativas normativas dos processos autônomos e máquinas. Vários órgãos e grupos, tal como o UK Engineering and Physical Scien-

ces Research Council (Conselho de Pesquisa de Engenharia e Ciências Físicas) (EPSRC), propuseram “princípios de robótica”, mas não existe um conjunto abrangente, ou global, de normas;

- IA e governança robótica: a falta de conhecimentos gerais sobre pesquisa em IA e aplicativos cria um desafio de previsão para as autoridades políticas. Além disso, é difícil determinar quais instituições devem tomar decisões sobre as políticas de IA. O reconhecimento desses fatores cria espaços para procedimentos inovadores de governança e a potencial criação de novos tipos de comissões, agências ou grupos consultivos, cuja autoridade deverá ainda ser cimentada;
- Resolução de conflitos: atualmente, não existem estruturas estabelecidas ou melhores práticas para resolver os conflitos associados aos sistemas e aplicações da IA. Várias dificuldades relacionadas à antecipação dos potenciais conflitos complicam o desenvolvimento dessas estruturas. Por exemplo, as pesquisas em IA não são regulamentadas – mesmo que os produtos que utilizam os aplicativos de IA o sejam – pondo, assim, o ônus regulatório no nível dos produtos.

O autor ainda segue mencionando que existem 10 questões que todo mundo deve saber sobre a IA: i) A IA muda ao longo do tempo; ii) A IA geral não existe, mas a IA estreita sim (Google, Siri, outras); iii) IA, robôs e seres humanos funcionam melhor quando trabalham juntos; iv) Os sistemas de IA precisam de orientação humana e o sucesso reside na IA ser treinada para observar pessoas e alinhar os atuais sistemas aos objetivos humanos; v) Ainda não se compreende como os sistemas de redes neurais artificiais funcionam, ou seja, alguns ainda são “caixas pretas”; vi) os recursos de IA são abertos e estão disponíveis; vii) para o uso da IA os dados dos indivíduos devem estar em ordem; viii) Até mesmo os sistemas mais inteligentes podem trazer respostas imprecisas ou preconceituosas; ix) A IA e robótica irão transformar as tarefas em vez de tornar os seres humanos obsoletos e x) a forma como os sistemas de IA são aplicados pelas organizações para os problemas da vida real é o principal fator de seu impacto, ou seja, na medida em que os sistemas forem se tor-

nando mais capazes, a tomada de decisão de onde e quando usa-los ganhará mais importância (Schwab, 2018, p. 189-192).

A Unesco (2021) emitiu assim em 2021 recomendações visando fornecer um marco universal de valores e princípios para os Estados, comunidade, empresas, atrelados a proteção dos Direitos Humanos e as liberdades fundamentais, bem como promover o diálogo e acesso equitativo aos avanços da IA com o compartilhamento de benefícios. Sendo listados os valores: i) Respeito, proteção, promoção dos direitos humanos e liberdades fundamentais e dignidade humana; ii) prosperidade ambiental e ecossistêmica; iii) garantir diversidade e inclusão; iv) viver em sociedades pacíficas, justas e interconectadas. Quanto aos princípios elenca: i)proporcionalidade e não causar dano; ii) segurança e proteção; iii) justiça e não discriminação; iv) sustentabilidade; v) direito a privacidade e proteção de dados; vi) supervisão humana e determinação; vii) transparência e explicabilidade; viii) responsabilidade e prestação de contas; ix) conscientização e alfabetização; x) governança e colaboração adaptáveis e com múltiplas partes interessadas.

Portanto, percebe-se que à IA, o Princípio da Precaução deve se revelar em ações éticas voltadas a proteção Humana, diante disso a Organização para Cooperação e Desenvolvimento Econômico (OCDE), tem trabalhado para determinar recomendações, tema este objeto do subtítulo a seguir.

#### **4. OS REFLEXOS DA PRECAUÇÃO COM OS RISCOS DA IA NAS ORGANIZAÇÕES INTERNACIONAIS**

Analisando-se os reflexos dos complexos efeitos emanados das novas tecnologias em um panorama internacional, atenta-se para recente manutenção aos princípios da Organização para Cooperação e Desenvolvimento Econômico (OCDE), alterada em 05 de fevereiro do corrente ano (OCDE, 2024), enaltecendo e incrementando o arcabouço principiológico tombado na *Recommendation of the Council on Artificial Intelligence* (OECD/LEGAL/0449), confeccionado em maio de 2019 (OCDE, 2019).

Conforme a recomendação do Conselho de Inteligência Artificial da OCDE, seus escopos são inspirados em atenção a Declaração Universal de Direitos Humanos de 1948, mas especialmente

na Resolução nº 70/1 da Assembleia Geral das Nações Unidas de 25 de setembro de 2015, que trata dos objetivos para o desenvolvimento sustentável, denominada *Transforming our World: The 2030 agenda for Sustainable Development* (OCDE, 2019). Os desígnios firmados pela Organização das Nações Unidas (ONU) na Resolução 70/1 de 2015 intentam na elaboração de um plano de ação, composto por 17 objetivos centrais e 169 metas benéficas aos cidadãos do mundo, visando alcançar cenários de prosperidade, paz e liberdade em um desenvolvimento sustentável que se ancora em três dimensões distintas, sejam a econômica, social e ambiental no ano de 2030 (ONU, 2015).

Cumpra-se frisar que os objetivos e metas estabelecidos na Resolução da Assembleia Geral das Nações Unidas seguem cinco eixos fundamentais para a agenda de desenvolvimento sustentável até 2030, três deles orbitando sob temas alinhados a boa tutela da natureza. O primeiro volta-se a intenção de gerar prosperidade globalizada, a fim de direcionar condutas para que todos os seres humanos tenham uma vida digna, próspera e plena por meio do progresso balanceado e harmônico entre a economia, a investida tecnológica e a natureza (ONU, 2015). Em sequência, compõe um dos eixos o tema do planeta, surgindo a preocupação transgeracional e a tomada de medidas para mitigar danos ao meio ambiente, observa-se:

Estamos determinados a proteger o planeta da degradação, nomeadamente através do consumo e da produção sustentáveis, da gestão sustentável dos seus recursos naturais e da tomada de medidas urgentes em matéria de alterações climáticas, para que possa apoiar as necessidades das gerações presentes e futuras (traduziu-se) (ONU, 2015).<sup>4</sup>

No terceiro eixo, encontra-se a preocupação com as pessoas, emergindo o intento em erradicar a pobreza e a fome, buscando confeccionar ações e metas que propiciem uma vida amparada por um ambiente saudável, em que possam exercer sua existência com dignidade e igualdade. Continuamente, o quarto eixo que surge na Resolução da ONU debruça-se sobre a paz, pretendendo incitar a promoção de sociedades justas, pacíficas e mais inclusivas, as

---

4 “We are determined to protect the planet from degradation, including through sustainable consumption and production, sustainably managing its natural resources and taking urgent action on climate change, so that it can support the needs of the present and future generations.” (ONU, 2015).

distanciando da violência, ressaltando-se no corpo do documento que não há desenvolvimento sustentável sem a promoção da paz, sequer paz, sem desenvolvimento sustentável. No mesmo sentido, sustenta-se na Resolução 70/1 a imperiosidade da edificação de uma solidariedade global para concretização da agenda de 2030 de desenvolvimento sustentável, centralizando-se nas necessidades das pessoas mais vulneráveis e na pobreza (ONU, 2015).

Entre os dezessete objetivos formulados na Resolução 70/1 da Assembleia Geral das Nações Unidas, destaca-se o i) décimo terceiro, descrevendo a necessidade de tomada de medidas urgentes para combater alterações climáticas e seus impactos, reforçar a capacidade de adaptação aos perigos relacionados com o clima e desastres naturais e melhorar a capacidade humana, sensibilização e educação humana voltada a mitigação dos riscos das alterações climáticas, adaptação, redução de impactos e alertas preditivos aos perigos. Ademais, tem-se o ii) décimo quarto objetivo volta-se para a conservação e utilização de forma sustentável dos recursos marinhos, mares e oceanos, tendo como uma das metas para o ano de 2025 a redução significativa da poluição marinha, especialmente fruto de atividades terrestres, de detritos à poluição por nutrientes, levando ao iii) décimo quinto objetivo busca proteger e promover o uso sustentável de ecossistemas terrestres, dirimir danos florestais, combater a degradação da terra e biodiversidade (ONU, 2015).

As preocupações acerca das incertezas da extensão e grau de dano à natureza e os seres humanos que emergem da Assembleia Geral da ONU têm seu influxo, como observou-se, nas diretrizes da Organização para Cooperação e Desenvolvimento Econômico (OCDE). Ao que tange os riscos emanados da Inteligência Artificial, o observatório da OCDE em IA reconhece que ainda que sistemas de IA sejam desenvolvidos com fins produtivos e benéficos, se utilizados de maneira errada, podem causar diversos malefícios as pessoas afetadas pela tecnologia, sendo fornecido como exemplo a reprodução de preconceitos pelos sistemas, tratando dos casos de discriminação algorítmica, a célere automação nos setores de produção, demandando o labor cada vez mais especializado, coleta invasiva de dados por equipamentos, produtos e serviços em Inteligência Artificial, concentração de poder dos desenvolvedores

e fornecedores de sistemas em IA, problemas em processos de tratamento de dados opacos e os efeitos incertos da tomada de decisão automatizada (OCDE, 2019).

No mesmo sentido, frisa-se que a Inteligência Artificial possui uma ampla gama de funcionalidades, e se tratando de uma ferramenta geral, a Organização para Cooperação e Desenvolvimento Econômico a observa como matéria de multidisciplinar e multilateral natureza, demandando um ativo papel cooperativo da comunidade internacional para definição de padrões adequados de desenvolvimento e difusão de soluções em IA (OCDE, 2019). Assim, a tônica disruptiva que surge do campo da Inteligência Artificial emerge como um desafio para problemas éticos-jurídicos, vez que o refinamento técnico acompanha formas mais sutis e menos perceptíveis dos possíveis riscos para o espaço vital das pessoas que podem ser impactadas (Sarlet; Caldeira, 2021, p. 173).

O empenho em tratar das incertezas que giram sob os efeitos do uso indiscriminado da Inteligência Artificial são visualizados nas diligências da Organização para Cooperação e Desenvolvimento Econômico desde meados de novembro de 2016, com a Conferência de Paris *Technology Foresight Forum on AI* (OCDE, 2016). Na oportunidade, Consultores de Inovação como Mr. Olivier Ezratty já alertavam sobre a imperiosidade de sensibilizar decisores políticos para investimentos em Inteligência Artificial e cuidados especiais em eixos estratégicos como transportes, educação universitária, saúde, meio ambiente e energia (OCDE, 2016).

Outro momento em que se vislumbra a preocupação com o célere desenvolvimento da nova tecnologia e amplo campo de aplicação é na Conferência *Intelligent Machines, Smart Policies in 2017*, em que participaram meia centena de palestrantes e mais de trezentos cientistas, motivados a dialogar não somente sobre as implicações da Inteligência Artificial no mundo real, mas da imperiosidade de estabelecer colaborações entre governo, indústria e políticas públicas para um desenvolvimento ético e aplicação segura da Inteligência Artificial em áreas como educação, segurança pública, vida pessoal e laboral (OCDE, 2018). No curso do evento o Conselheiro do Departamento de Indústria, Inovação e Ciência da Austrália, Alexander Cooke, apresentou o DEA (*Digital Earth Australia*), exemplificando como a IA pode ser manuseada para monitorar a irrigação ambiental, modelamento de

planos para bacias, assim como formular aplicabilidades para questões de dinâmica ecológica, observa-se:

Alexander Cooke, Conselheiro do Departamento de Indústria, Inovação e Ciência da Austrália (Digital Earth Australia) apresentou o projeto Digital Earth Australia (DEA) estabelecido pelo governo australiano para aumentar a eficiência e eficácia de programas e políticas públicas que precisam de informações precisas e oportunas informações espaciais sobre a saúde e produtividade da paisagem da Austrália. A DEA utiliza dados do Sentinel 2, que são recolhidos e analisados com técnicas de aprendizagem automática para modelar planos de bacias e monitorizar e modelar a irrigação ambiental, o oceano, as costas e para fornecer respostas ecológicas dinâmicas (traduziu-se) (OCDE, 2018).<sup>5</sup>

O arcabouço de contribuições científicas, acadêmicas, tanto do setor público quanto privado observadas nas conferências da Organização para Cooperação e Desenvolvimento Econômico desaguará na criação do *Council Recommendation on Artificial Intelligence* (OCDE, 2019), almejando trazer recomendações em nível internacional para edificação de sistemas de Inteligência Artificial seguros e confiáveis em diversos eixos de aplicação, seguindo uma linha principiológica consolidada em cinco princípios centrais, quais foram arrolados na Recomendação do Conselho de Inteligência Artificial da OCDE, aprovado em 21 de maio de 2019 (Morandin-Ahuerma, 2023, p. 104). Atentando-se aos ditames da recomendação internacional, a OCDE apresenta diretrizes com valores centralizados na figura humana, em sua dignidade, na necessidade de transparência, segurança, estabilidade dos sistemas de IA, desenvolvimento inclusivo, sustentável, visando uma contribuição à comunidade global nos movimentos exploratórios desta tecnologia (Tuma; Mathias Júnior; Penha. 2023, p. 45).

---

5 “Alexander Cooke, Counsellor, Department of Industry, Innovation and Science, Australia (Digital Earth Australia) presented the Digital Earth Australia (DEA) project established by the Australian government to increase the efficiency and effectiveness of public programmes and policies that need accurate and timely spatial information on the health and productivity of Australia’s landscape. DEA uses Sentinel 2 data, which are collected and analysed with machine learning techniques to model basin plans and monitor and model environmental watering, ocean, coasts and to provide dynamic ecological responses.” (OCDE, 2018).

Analisando-se os princípios firmados pela OCDE, observa-se o i) princípio do crescimento inclusivo, desenvolvimento sustentável e bem-estar, almejando edificar um panorama que vise angariar resultados positivos para a população internacional e inclusive para o planeta, visando produzir efeitos benéficos para tutela de ambientes naturais e sustentabilidade ambiental. Ainda voltado aos riscos da IA, tem-se o ii) princípio do respeito ao Estado de Direito, direitos humanos, valores democráticos, justiça e privacidade, importando analisar-se a alínea “b)” da recomendação, que expressa a imperiosidade dos intervenientes criarem mecanismos para salvaguarda de riscos de desvio de finalidade de atividade de sistemas de Inteligência Artificial, assim como do seu uso indevido, intencional ou não, conjuntamente com a necessidade de supervisionamento humano (OCDE, 2024).

O terceiro é o iii) princípio da transparência, para promoção de uma compreensão acerca das capacidades e limitações dos sistemas de IA, fornecendo inclusive informações a pessoas afetadas negativamente por soluções em Inteligência Artificial. Sucessivamente, o iv) princípio da robustez, que orbita sob a tônica de segurança necessária na construção de um sistema de IA, ressaltando-se a importância de sedimentar formas de mitigar, reduzir ou impedir impactos indevidos pela Inteligência Artificial (OCDE, 2024). Como um desdobramento deste princípio, verifica-se a busca por dissolver a nebulosidade dos efeitos negativos que um sistema de IA posa produzir, imperando a rastreabilidade da Inteligência Artificial, observa-se:

Garantizar la trazabilidad es esencial para promover la responsabilidad y la transparencia en los sistemas de IA. Mediante el seguimiento y el análisis de todos los conjuntos de datos, procesos y decisiones pertinentes, los involucrados pueden contribuir a garantizar que los sistemas de IA funcionen de forma ética y responsable, en consonancia con las mejores prácticas y tecnologías (Morandin-Ahuerma, 2023, p. 107).

Por fim, o último componente do arcabouço principiológico da Recomendação do Conselho de Inteligência Artificial da OCDE é o v) Princípio da Responsabilidade, vinculando a observância aos quatro primeiros princípios no desenvolvimento responsável de sistemas de Inteligência Artificial, assegurando que suas diretrizes

caminhem rumo à resultados éticos e confiáveis. Neste contexto, os desenvolvedores devem atentar-se nas questões acerca da rastreabilidade do sistema durante todo seu ciclo de vida, recomendando-se a adoção de abordagens de tutela baseadas na gestão de riscos emanados pelo sistema de Inteligência Artificial (OCDE, 2024).

Consoante as premissas emanadas pela própria Organização para Cooperação e Desenvolvimento Econômico, no contexto de impacto climático, os princípios firmados pela OCDE em termos de ação, dissolvem as imediatas promessas utópicas da nova tecnologia, rejeitando a via do solucionismo tecnológico (OCDE, 2021). Assim, observa-se um tombamento principiológico que não subestima os riscos e incertezas acerca dos efeitos da IA no contexto ambiental, mas trata de promover auxílios para movimentações políticas responsáveis, assim como as acadêmicas e de iniciativa da comunidade civil. Analisa-se:

Embora tanto as alterações climáticas como a IA apresentem incertezas profundas – até mesmo riscos existenciais – a nossa situação ambiental autocriada exige-nos que atuemos ainda mais rápida e que utilizemos todas as ferramentas à nossa disposição (traduziu-se) (Dennis; Kirnberger; Shankar, 2023).<sup>6</sup>

No tocante, vislumbra-se nas diretrizes do *Council Recommendation on Artificial Intelligence* relevante norte para mitigação dos riscos incertos que emanam da Inteligência Artificial, seus princípios firmados pela OCDE em 2019 e ratificados no percurso de 2024 tem reflexos em organizações de desenvolvimento de normas para produtos e serviços como a Organização Internacional de Padronização (ISO), que harmonicamente com os princípios do crescimento inclusivo, bem-estar e desenvolvimento sustentável, declarou que as soluções em Inteligência Artificial não devem guiar o futuro a um estado de ameaça a vida humana, a sua saúde, aos seus bens e especialmente ao meio ambiente (Dennis; Kirnberger; Shankar, 2023).

Outro reflexo da necessidade de manter a precaução com os impactos a vida e ao planeta pode ser verificado no Acordo Administrativo *U.S – EU Trade and Technology Council Commitment*, firmado

---

6 “While both climate change and AI present deep uncertainties – even existential risks – our self-created environmental predicament demands us to act even faster and to utilize all tools at our disposal.” (OCDE, 2023).

entre Estados Unidos da América (EUA) e União Europeia (UE) em janeiro de 2023 (EUA, 2023). Conforme o acordo internacional, reúnem-se especialistas tanto dos EUA quanto União Europeia para enfrentar os principais desafios globais ligados a Inteligência Artificial, firmando cinco focos para o Conselho, sejam i) previsões climáticas extremas, ii) estipulação de uma gestão de resposta a casos emergenciais, iii) refinamento de aplicações na área da saúde, iv) melhoramento de aplicações em rede elétrica e v) otimização na área da agricultura (EUA, 2023).

Frontalmente ao que se observa, os princípios firmados no corpo da Recomendação do Conselho de Inteligência Artificial da OCDE compõem imperioso norte ético para trilhar caminho rumo à uma aplicação segura da Inteligência Artificial (Sarlet; Caldeira. 2021, p. 170). Assim, constata-se que muito embora a tecnologia emergente apresente inovadoras aplicações e funcionalidades, ela traz consigo novos desafios, riscos e efeitos incertos, quais, em caso de contínua exploração desregrada e desvencilhada de balizas éticas, podem impactar não apenas nossa geração atual, mas todas as que se sucederão. Neste sentido, diligenciam organizações internacionais como a Organização para Cooperação e Desenvolvimento Econômico em sentido à uma precaução frente aos possíveis impactos da Inteligência Artificial, trazendo recomendações a decisores e gestores políticos, desenvolvedores de sistemas e a comunidade internacional.

## **CONSIDERAÇÕES FINAIS**

No presente trabalho verificou-se que os sistemas de Inteligência Artificial que atualmente fazem parte do dia-a-dia de inúmeras pessoas ao redor do mundo, numa pluralidade de aplicações e ações algorítmicas que se relacionam a quase todos os setores da vida em sociedade e que inúmeros esforços teóricos, acadêmicos e científicos acerca da possibilidade de uma máquina simular a capacidade cognitiva humana estão em voga. Ao mesmo tempo, verificou-se que as inovações tecnológicas jamais devem vir desacompanhadas da sua orientação ética, haja vista que, ao mesmo tempo que proporciona inúmeras oportunidades trazem consigo riscos de graves danos à Humanidade.

Diante disso, chama-se o Princípio da Precaução como instrumento para cautela das IAs e percebe-se que diferentemente da aplicação da precaução quanto as outras novas tecnologias, pois em sua maioria a avaliação de riscos é voltada para os resultados em pesquisas científicas, nas ações das IAs a precaução circunda mais em medidas éticas, protetivas aos Direitos Humanos (vida, liberdade, honra, privacidade, emprego...).

Concluindo-se, assim, que apesar da IA permitir soluções inovadoras ela também traz consigo novos desafios, riscos e resultados incertos, ou seja, abstralidades ainda não desenhadas por completo pelas orientações éticas. Assim, na medida do seu desenvolvimento ela deve vir balizada pela ética para respeitar os Direitos Humanos em sua completude para as presentes e futuras gerações. Para tanto, compreende-se que são pertinentes as recomendações desenhadas pelas organizações internacionais como a Organização para Cooperação e Desenvolvimento Econômico em seu do Conselho de Inteligência Artificial visando a precaução diante dos riscos da Inteligência Artificial.

## **REFERÊNCIAS**

- ANTISERI, Dario; REALE, Giovanni. *Filosofia: Idade Moderna*. Porto Alegre: Paulus, 2018.
- BECK, Ulrich. *A metamorfose do mundo: novos conceitos para uma nova realidade*. Rio de Janeiro: Zahar, 2018.
- BERENTE, Nicholas *et al.* A. Managing Artificial Intelligence. *MIS Quarterly*, v. 45, n. 3, p. 1433-1450, 2021.
- BERWIG, Juliane Altmann; ENGELMANN, Wilson. O princípio da precaução: diretrizes para sua aplicação empírica. *Revista Pensar*, v. 28, n.1, p.1-14, 2023.
- DENNIS, Claire; KIRNBERGER, Johannes Leon; SHANKAR, Var. We need to use AI to fight climate change. Paris: OCDE, 30 maio 2023. Disponível em: <https://oecd.ai/en/wonk/fight-climate-change>. Acesso em: 12 de jun. de 2024.

- ESTADOS UNIDOS DA AMÉRICA (EUA). Casa Branca. *Statement by National Security Advisor Jake Sullivan on the New U.S.-EU Artificial Intelligence Collaboration*. Washington: EUA, 2023. Disponível em: <https://www.whitehouse.gov/briefing-room/statements-releases/2023/01/27/statement-by-national-security-advisor-jake-sullivan-on-the-new-u-s-eu-artificial-intelligence-collaboration>. Acesso em: 12 jun. 2024
- HAN, Byung-Chul. *Sociedade paliativa: A dor hoje*. Petrópolis: Vozes, 2021.
- KAUFMAN, Dora. Inteligência artificial: repensando a mediação. *Brazilian Journal of Development*, v. 6, n. 9, p. 1-19, 2020. Disponível em: <https://ojs.brazilianjournals.com.br/ojs/index.php/BRJD/article/view/16481>. Acesso em: 11 jun. 2024.
- LEE, Kai-fu. *Inteligência Artificial: Como os robôs estão mudando o mundo, a forma como nos amamos, nos comunicamos e vivemos*. Rio de Janeiro: Globo Livros, 2019.
- LÉVY, Pierre. IEML: rumo a uma mudança de paradigma na Inteligência Artificial. *Matrizes*, v. 16, n. 1, p. 11-34, jan. 2022. Disponível em: <https://www.redalyc.org/journal/1430/143071289015/143071289015.pdf>. Acesso em: 11 jun. 2024.
- OLIVEIRA, Jelson. *Compreender Hans Jonas*. Petrópolis: Vozes, 2014.
- ORGANIZAÇÃO DAS NAÇÕES UNIDAS (ONU). *Resolution adopted by the General Assembly on 25 September 2015: Transforming our world, the 2030 Agenda for Sustainable Development*. Nova Iorque: Assembleia Geral da ONU, 2015. Disponível em: <https://documents.un.org/doc/undoc/gen/n15/291/89/pdf/n1529189.pdf>. Acesso em: 11 jun. 2024.
- ORGANIZAÇÃO DAS NAÇÕES UNIDAS PARA A EDUCAÇÃO, A CIÊNCIA E A CULTURA (UNESCO). *Recomendação sobre a Ética da Inteligência Artificial*. Paris: Unesco, 2021. Disponível em: [https://unesdoc.unesco.org/ark:/48223/pf0000381137\\_por](https://unesdoc.unesco.org/ark:/48223/pf0000381137_por). Acesso em: 31 jul. 2024.
- ORGANIZAÇÃO PARA A COOPERAÇÃO E DESENVOLVIMENTO ECONÔMICO (OCDE). *Recommendation of the Council on Artificial Intelligence*. OECD/LEGAL/0449. Paris: OCDE, maio 2019. Disponível em: <https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0449>. Acesso em: 11 jun. 2024.

- ORGANIZAÇÃO PARA A COOPERAÇÃO E DESENVOLVIMENTO ECONÔMICO (OCDE). *Recommendation of the Council on Artificial Intelligence*. OECD/LEGAL/0449. Paris: OCDE, maio 2024. <https://legalinstruments.oecd.org/en/instruments/oecd-legal-0449>. Acesso em: 11 jun. 2024.
- ORGANIZAÇÃO PARA A COOPERAÇÃO E DESENVOLVIMENTO ECONÔMICO (OCDE). *The Context*. Paris: OECD AI Policy Observatory, [2020]. Disponível em: <https://oecd.ai/en/about/the-context>. Acesso em: 11 jun. 2024.
- ORGANIZAÇÃO PARA A COOPERAÇÃO E DESENVOLVIMENTO ECONÔMICO (OCDE). *Technology Foresight Forum 2016 on Artificial Intelligence (AI)*. Paris: OCDE, 2016. Disponível em: <https://www.oecd.org/sti/ieconomy/technology-foresight-forum-2016.htm>. Acesso em: 11 jun. 2024
- ORGANIZAÇÃO PARA A COOPERAÇÃO E DESENVOLVIMENTO ECONÔMICO (OCDE). *AI, Intelligent Machines, Smart Policies Conference Summary*. DSTI/CDEP(2018)8/FINAL. Paris: OCDE, 2 ago. 2018. Disponível em: <https://www.oecd-ilibrary.org/docserver/f1a650d9-en.pdf>. Acesso em: 12 jun. 2024.
- SARLET, Gabrielle Bezerra Sales; CALDEIRA, CRISTINA. A Inteligência Artificial e o ecossistema industrial na sua relação com as patentes na área da saúde: Uma abordagem antropocêntrica sobre os desafios impostos em tempos de pandemia. *Privacy and Data Protection Magazine*, n. 1, p. 154-199, 2021. Disponível em: [https://repositorio.pucrs.br/dspace/bitstream/10923/18749/2/A\\_INTELIGENCIA\\_ARTIFICIAL\\_E\\_O\\_ECOSSITEMA\\_INDUSTRIAL\\_NA\\_SUA\\_RELACAO\\_COM\\_AS\\_PATENTES\\_NA\\_AREA\\_DA\\_SADE\\_UMA\\_ABORDAGEM.pdf](https://repositorio.pucrs.br/dspace/bitstream/10923/18749/2/A_INTELIGENCIA_ARTIFICIAL_E_O_ECOSSITEMA_INDUSTRIAL_NA_SUA_RELACAO_COM_AS_PATENTES_NA_AREA_DA_SADE_UMA_ABORDAGEM.pdf). Acesso em: 11 jun. 2024.
- SCHWAB, Klaus. *A Quarta Revolução Industrial*. São Paulo: Edipro, 2016.
- SILVEIRA, Paulo Antônio Caliendo Velloso da. *Ética e Inteligência Artificial: da possibilidade filosófica de Agentes Morais Artificial*. Porto Alegre: Fi, 2021.
- SUNSTEIN, Cass R. *Laws of fear: beyond the precautionary principle*. New York: Cambridge, 2005.

- TAULLI, Tom. *Introdução à Inteligência Artificial: uma abordagem não técnica*. São Paulo: Novatec, 2020.
- TUMA, Eduardo; MATHIAS JUNIOR, Alexandre Ferreira Mathias; PENHA, Thaluana Alves da. Acessibilidade e mobilidade do deficiente visual: contribuições no âmbito das cidades inteligentes o caso da Vila Clementino. *Revista Controle-Doutrina e Artigos*, v. 21, n. 2, p. 28-62, 2023. Disponível em: <https://revistacontrole.tce.ce.gov.br/index.php/RCDA/article/view/835>. Acesso em: 12 jun. 2024.