

**UNIVERSIDADE DO VALE DO ITAJAÍ – UNIVALI
VICE-REITORIA DE PESQUISA, PÓS-GRADUAÇÃO E INOVAÇÃO
PROGRAMA DE PÓS-GRADUAÇÃO STRICTO SENSU EM CIÊNCIA JURÍDICA– PPCJ
CURSO DE MESTRADO EM CIÊNCIA JURÍDICA – CMCJ
ÁREA DE CONCENTRAÇÃO: FUNDAMENTOS DO DIREITO POSITIVO
LINHA DE PESQUISA: DIREITO, JURISDIÇÃO E INTELIGÊNCIA ARTIFICIAL**

**UM ESTUDO SOBRE A CONSTRUÇÃO DA INTELIGÊNCIA
ARTIFICIAL DE CONFIANÇA SOB O ENFOQUE DOS DIREITOS
HUMANOS**

ANDRESA SILVEIRA ESTEVES

Itajaí-SC, fevereiro de 2022

UNIVERSIDADE DO VALE DO ITAJAÍ – UNIVALI
VICE-REITORIA DE PESQUISA, PÓS-GRADUAÇÃO E INOVAÇÃO
PROGRAMA DE PÓS-GRADUAÇÃO STRICTO SENSU EM CIÊNCIA JURÍDICA-PPCJ
CURSO DE MESTRADO EM CIÊNCIA JURÍDICA – CM CJ
ÁREA DE CONCENTRAÇÃO: FUNDAMENTOS DO DIREITO POSITIVO
LINHA DE PESQUISA: DIREITO, JURISDIÇÃO E INTELIGÊNCIA ARTIFICIAL

**UM ESTUDO SOBRE A CONSTRUÇÃO DA INTELIGÊNCIA
ARTIFICIAL DE CONFIANÇA SOB O ENFOQUE DOS DIREITOS
HUMANOS**

ANDRESA SILVEIRA ESTEVES

Dissertação submetida ao Curso de Mestrado em Ciência Jurídica da Universidade do Vale do Itajaí –UNIVALI, como requisito parcial à obtenção do título de Mestre em Ciência Jurídica.

Orientadora: Profa. Doutora Dirajaia Esse Pruner

Coorientador: Prof. Doutor Alexandre Moraes da Rosa

Itajaí-SC, fevereiro de 2022

DEDICATÓRIA

Dedico este trabalho de dissertação a minha mãe Rejane Araújo Silveira (in memoriam) por me ensinar que tudo segue o caminho e o tempo certo para se realizar, e que não existe transformação sem desafios, sem aceitação e autoconhecimento, a resiliência é crucial para alcançarmos os nossos maiores sonhos. Dedico também ao meu irmão Allan Luã Soares, pois, como em todos os momentos de minha vida, sempre foi meu maior companheiro, minha força e minha fonte de confiança, você traz o colorido e a leveza que fazem tudo valer a pena.

AGRADECIMENTOS

Gratidão! Essa é a palavra que representa esta etapa da minha vida, e principalmente, por ter a sorte de poder contar com tantas pessoas merecedoras dessa gratidão. Porém, por mais que possa parecer estranho ou egocêntrico, diante de todos os desafios que a vida trouxe no decorrer desta conquista, que foi tão almejada por mim, e que por muitas vezes cheguei a acreditar que não se realizaria, agradeço por sempre ter encontrado motivos para seguir em frente, por ter sido persistente e resiliente. Desistir se mostrou como o caminho por diversas vezes, mas nunca se tornou a opção e, por isso, eu sou minha maior fonte de gratidão.

No entanto, não somos nada sozinhos, somos o somatório da energia depositadas pelos que nos rodeiam. Sendo assim, primeiramente, gostaria de agradecer meu irmão Allan Luã Soares, minha cunhada/irmã/melhor amiga Giovanna Richert Bahia Bittencout, minha irmã Moara Brito e minha melhor amiga/irmã/comadre Juliana Martins, por nunca me deixarem desistir, por despertarem em mim a minha melhor versão, vocês foram combustíveis de força e confiança, tornaram esse processo possível com uma lealdade incontestável.

Agradeço a minha mãe Salete Silveira Esteves (sim, eu tive a sorte de ter duas mães) e ao meu pai Gelson Galeno dos Santos Esteves (in memoriam) por terem sido a fonte de vida, por terem sido a base que sustenta toda a minha construção.

Aos meus irmãos, Caroline Silveira Esteves, Bruno Crivelatti e Marum Brito, integrantes de um time de 8 irmãos, esse time grande que está sempre pronto tudo, vocês são minha maior alegria.

Gostaria de demonstrar minha gratidão, de uma forma muito especial, ao meu orientador Dr. Alexandre Moraes da Rosa, por ter apostado suas fichas em mim desde o início, por todas as oportunidades que me proporcionou, que são inúmeras e por todos os ensinamentos, nunca saberia como demonstrar claramente o tamanho da gratidão que carrego comigo.

Agradeço a minha orientadora Dra. Dirajaia Esse Pruner, por ter enfrentado de coração aberto o grande desafio de me orientar na construção desta pesquisa, você foi incrível e o um dos maiores presentes que a Univali me deu, gratidão pela paciência, pela confiança, pelo carinho, e, principalmente, por toda contribuição que foi essencial nessa trajetória.

Quero agradecer, também de uma forma especial, a secretária da Univali,

Cristina de Oliveira Goncalves Koch, por toda compreensão, pela paciência e, principalmente, por ser sempre tão prestativa e atenciosa.

Ainda, gostaria de agradecer a todo corpo docente da Univali, em especial, aos professores, Dra. Denise Schmitt Siqueira Garcia, Dra. Heloise Siqueira Garcia, Dra, Emmanoela Cristina Andrade Lacerda, Dr. Pedro Manoel Abreu e o Dr. José Antônio Savaris, por enriquecerem meu senso crítico com tantos ensinamentos.

Por fim, meu agradecimento mais que especial vai aos meus filhos pets, Zoltan e Pretinha, por serem minhas melhores companhias, a alegria dos meus dias e a fonte de amor mais pura e inesgotável que já vivi.

TERMO DE ISENÇÃO DE RESPONSABILIDADE

Declaro, para todos os fins de direito, que assumo total responsabilidade pelo aporte ideológico conferido ao presente trabalho, isentando a Universidade do Vale do Itajaí, a Coordenação do Curso de Mestrado em Ciência Jurídica, a Banca Examinadora e o Orientador de toda e qualquer responsabilidade acerca dele.

Itajaí, fevereiro de 2022.

A handwritten signature in blue ink, consisting of stylized, overlapping loops and strokes, representing the name Andresa Silveira Esteves.

Andresa Silveira Esteves

Mestranda

PÁGINA DE APROVAÇÃO

MESTRADO

Conforme Ata da Banca de Defesa de Mestrado, arquivada na Secretaria do Programa de Pós-Graduação *Stricto Sensu* em Ciência Jurídica PPCJ/UNIVALI, em 17/03/2022, às 14 horas, a mestranda ANDRESA SILVEIRA ESTEVES fez a apresentação e defesa da Dissertação, sob o título “UM ESTUDO SOBRE A CONSTRUÇÃO DA INTELIGÊNCIA ARTIFICIAL DE CONFIANÇA SOB O ENFOQUE DOS DIREITOS HUMANOS”.

A Banca Examinadora foi composta pelos seguintes professores: Doutora Dirajaia Esse Pruner (UNIVALI) como presidente e orientadora, Doutora Cleide Calgaro (UCS) como membro, Doutor Marcos Leite Garcia (UNIVALI) como membro e Doutor Alexandre Moraes Da Rosa (UNIVALI) como membro suplente. Conforme consta em Ata, após a avaliação dos membros da Banca, a Dissertação foi Aprovada.

Por ser verdade, firmo a presente.

Itajaí (SC), 17 de março de 2022.



PROF. DR. PAULO MÁRCIO DA CRUZ
Coordenador/PPCJ/UNIVALI

ROL DE ABREVIATURAS E SIGLAS

AGNU – Assembleia Geral das Nações Unidas

CNJ – Conselho Nacional de Justiça

COMPAS – Correctional Offender Management Profiling for Alternative Sanctions

DUDH – Declaração Universal de Direitos Humanos

EBIA – Estratégia Brasileira de Inteligência Artificial

GPAN IA – Grupo de Peritos de Alto Nível sobre a Inteligência Artificial

IA – Inteligência Artificial

LGPD – Lei Geral de Proteção de Dados

OCDE – Organização de Cooperação pelo Desenvolvimento Econômico

ONU – Organização das Nações Unidas

PL – Projeto de Lei

PLN – Processamento de Linguagem Natural

RNA – Rede Neural Artificial

STF – Supremo Tribunal Federal

STJ – Superior Tribunal de Justiça

TCU – Tribunal de Contas da União

RESUMO

A presente dissertação está inserida na linha de pesquisa “Direito, Jurisdição e Inteligência Artificial” e teve como objetivo geral analisar os parâmetros necessários para que a construção de uma inteligência artificial de confiança esteja inserida em parâmetros éticos, bem como nos Direitos Humanos e os Direitos Fundamentais. Sendo assim, a definição de um Marco Legal sobre inteligência artificial ainda é um desafio a ser enfrentado. O problema de pesquisa que se buscou responder é sobre quais os parâmetros que a IA de Confiança deve adotar para o respeito aos Direitos Humanos e os Direitos Fundamentais? O método utilizado para a abordagem foi o hipotético-dedutivo. Quanto a abordagem, adotou-se o tipo qualitativo e, no que se refere a técnica de pesquisa, utilizou-se a pesquisa bibliográfica, legislação artigos científicos em meios eletrônicos e notícias em sites oficiais. Para o desenvolvimento da pesquisa o trabalho foi dividido em três capítulos. No primeiro capítulo, se apresentou brevemente a evolução histórica, bem como algumas conceituações essenciais para a compreensão do desenvolvimento da inteligência artificial. No segundo capítulo, se estudou alguns sistemas de inteligência artificial já implementados no sistema judiciário tanto no âmbito internacional como no âmbito nacional, ainda neste capítulo se analisou as propostas de regulamentação sobre inteligência artificial no mundo e no Brasil. No que se refere ao terceiro capítulo, foi apresentado o conceito da Teoria Crítica dos Direitos Humanos, apresentado por Herrera Flores, a qual busca contradizer a ideia de universalidade dos Direitos Humanos, propondo que se reconheça a complexidade e pluralidade existe na sociedade. Por fim, se buscou analisar os parâmetros dispostos nas Orientações Éticas para uma IA de Confiança apresentados pela Comissão Europeia, e a partir deste documento construir parâmetros para a regulamentação brasileira. É importante salientar, que o que se pretendeu é uma profunda reflexão acerca das implicações éticas, legais e sociais que a implementação da inteligência artificiais causam na sociedade. Com isso, para ser possível a instalação de sistemas de inteligência artificial no poder judiciário brasileiro, assim como no âmbito privado, se faz essencial que haja uma confiança entre aqueles que estarão sujeitos a máquina.

Palavras-chaves: Inteligência Artificial de Confiança. Direitos Humanos. Ética.

ABSTRACT

This dissertation is part of the line of research “Law, Jurisdiction and Artificial Intelligence”. Its general objective is to analyze the parameters needed to construct a reliable artificial intelligence to be inserted in ethical parameters, as well as in Human Rights and Fundamental Rights. The definition of a Legal Framework on artificial intelligence is still a challenge to be faced. The research problem that we sought to answer is: which parameters should the AI of Trust adopt to respect Human Rights and Fundamental Rights? The method used for the approach was the hypothetical-deductive. A qualitative approach was adopted, and for the research technique, a bibliographic review was carried out, which included legislation, scientific articles in electronic media, and news on official websites. This work is divided into three chapters. The first chapter presents the historical evolution of artificial intelligence, as well as some essential concepts for understanding its development. The second chapter examines some artificial intelligence systems already implemented in the judiciary both internationally and nationally, and analyzes the proposals for regulation on artificial intelligence in the world and in Brazil. The third chapter presents Herrera Flores’ concept of Critical Theory of Human Rights, which seeks to contradict the idea of universality of Human Rights, proposing that the complexity and plurality that exist in society be recognized. Finally, this work analyzes the parameters set out in the Ethical Guidelines for a Trustworthy AI presented by the European Commission, and from this document, seeks to build parameters for the Brazilian regulation. The ultimate goal of this study is to offer a deep reflection on the ethical, legal and social implications of the implementation of artificial intelligence for society. In order to be able to install artificial intelligence systems in the Brazilian judiciary, as well as in the private sphere, it is essential for there to be trust between those who will be subject to the machine.

Keywords: Reliable Artificial Intelligence. Human rights. Ethic.

SUMÁRIO

INTRODUÇÃO	12
1 COMPREENDENDO O APRENDIZADO DE MÁQUINA COM USO DA INTELIGÊNCIA ARTIFICIAL	15
1.1 BREVE EVOLUÇÃO HISTÓRICA DA INTELIGÊNCIA ARTIFICIAL E ALGUNS CONCEITOS RELACIONADOS AO APRENDIZADO DE MÁQUINA	15
1.1.1 Breve evolução histórica da inteligência artificial	15
1.1.2 Conceito de inteligência artificial	18
1.1.3 Big Data.....	21
1.1.4 Algoritmo	24
1.1.5 Black Box	26
1.1.6 Processamento de Linguagem Natural (PLN)	27
1.2 INTELIGÊNCIA ARTIFICIAL FRACA, FORTE E SUPERINTELIGÊNCIA.....	29
1.2.1 Inteligência Artificial Fraca.....	29
1.2.2 Inteligência Artificial Forte	30
1.2.3 Inteligência Artificial Superinteligente.....	31
1.3 TÉCNICAS DE APRENDIZAGEM EM INTELIGÊNCIA ARTIFICIAL	33
1.3.1 Machine learning	33
1.3.1.1 Aprendizagem supervisionada	36
1.3.1.2 Aprendizagem não supervisionada	38
1.3.1.3 Aprendizagem por reforço	40
1.3.2 Deep Learning	42
1.3.3 Redes Neurais Artificiais	45
2 A APLICAÇÃO DA INTELIGÊNCIA ARTIFICIAL NO DIREITO E A SUA ATUAL REGULAMENTAÇÃO	48
2.1 PRINCIPAIS FERRAMENTAS DA INTELIGÊNCIA ARTIFICIAL UTILIZADAS NO ÂMBITO JURÍDICO EM PÂRAMETRO INTERNACIONAL	48
2.1.1 Sistema COMPAS	49
2.1.2 Robô ROSS.....	51
2.1.3 LawGeex	53
2.2 PRINCIPAIS FERRAMENTAS DA INTELIGÊNCIA ARTIFICIAL UTILIZADAS NO ÂMBITO JURÍDICO EM PÂRAMETRO NACIONAL.....	54
2.2.1 Sistema Victor	54

2.2.2 Projeto Sócrates.....	57
2.2.3 Plataforma de Inteligência Artificial Dra. Luzia	59
2.2.4 Sistemas Alice, Sofia e Mônica	60
2.2.4.1 Alice.....	61
2.2.4.2 Sofia	63
2.2.4.3 Monica.....	64
2.3 REGULAMENTAÇÃO DA INTELIGÊNCIA ARTIFICIAL NO MUNDO E NO BRASIL	65
2.3.1 As propostas para regulamentar a inteligência artificial ao redor do mundo	66
2.3.2 Passos em direção da regulamentação sobre o uso da Inteligência Artificial no Brasil	71
3 UM ESTUDO SOBRE A CONSTRUÇÃO DA INTELIGÊNCIA ARTIFICIAL DE CONFIANÇA SOB O ENFOQUE DOS DIREITOS HUMANOS.....	76
3.1 O PROCESSO DE INTERNACIONALIZAÇÃO E UNIVERSALIZAÇÃO DOS DIREITOS HUMANOS	76
3.1.1 O caráter universal da Declaração Universal de Direitos Humanos de 1948 ...	77
3.1.2 Teoria crítica dos Direitos Humanos sob o olhar de Herrera Flores	80
3.2 ÉTICA NO ÂMBITO JURÍDICO	89
3.3 PARÂMETROS PARA REGULAMENTAÇÃO BRASILEIRA SOBRE O USO DA INTELIGÊNCIA ALRTIFICIAL COM BASE NAS ORIENTAÇÕES ÉTICAS DA COMISSÃO EUROPÉIA	94
3.3.1 Componentes de uma Inteligência Artificial de confiança	97
3.3.2 Princípios éticos bases de uma Inteligência Artificial de Confiança	103
3.3.3 Requisitos essenciais para a concretização de uma Inteligência Artificial de confiança.....	108
CONCLUSÃO	126
REFERÊNCIAS.....	133

INTRODUÇÃO

A presente dissertação está inserida na linha de pesquisa “Direito, Jurisdição e Inteligência Artificial” e tem como objetivo institucional a obtenção do título de Mestre em Ciência Jurídica pelo Curso de Mestrado em Ciência Jurídica da Universidade do Vale do Itajaí – UNIVALI.

Tem como tema central a regulamentação da Inteligência Artificial de Confiança sob o enfoque dos Direitos Humanos. Com isso, tem o objetivo geral de analisar os parâmetros necessários para que a construção dessa regulamentação esteja inserida em parâmetros éticos, bem como nos Direitos Humanos e os Direitos Fundamentais. Os objetivos específicos consistem em:

- Discorrer brevemente sobre a evolução histórica da inteligência artificial e explicar alguns conceitos essenciais relacionados a aprendizagem de máquina.
- Apresentar algumas ferramentas de inteligência artificial utilizadas no âmbito jurídico, bem como a regulamentação existente tanto em âmbito internacional, como nacional.
- Analisar os parâmetros necessários para a construção da inteligência artificial de confiança a partir dos conceitos de Direitos Humanos e da Ética.

Salienta-se que a presente pesquisa tem relevância e se justifica, pois, tendo em vista que é notório que a tecnologia se faz cada vez mais presente no cotidiano das pessoas. Somos rodeados por computadores, tablets, celulares; passamos o dia conectados com outros indivíduos por meio da internet e das mais diversas redes sociais, as assistentes virtuais já assumem o controle das nossas atividades do dia, bem como o próprio judiciário já vem implementando robôs com uso da inteligência artificial visando otimizar o serviço. Além das funcionalidades já citadas, a tecnologia também pode fazer escolhas em nosso lugar, os algoritmos já são suficientemente desenvolvidos para tomarem, autonomamente, decisões baseadas em um banco de dados são fornecidas a eles.

Contudo, apesar da tecnologia de inteligência artificial já estar inserida na sociedade, ela vem avançando cada vez mais em uma velocidade imensurável, no entanto, ainda existe uma carência em relação a regulamentação dessa tecnologia. Num contexto mundial, alguns Países já avançaram em direção ao Marco Legal da Inteligência Artificial, dentre os principais estão a China, o Brasil, Estados Unidos e União Europeia.

No entanto, ao se debater sobre a construção de uma IA de Confiança, deve-se levar em consideração a construção de uma governança ética, garantindo a mitigação de consequências danosas às pessoas no uso dessa tecnologia, garantindo o respeito aos Direitos Humanos e os Direitos Fundamentais. Sendo assim, a construção de normas técnicas e legais, para instituir padrões de segurança e ética a serem seguidos no desenvolvimento de Sistemas de Inteligência Artificial, visa que a tecnologia além de benéfica, seja segura e ética.

Sendo assim, a definição de um Marco Legal sobre inteligência artificial ainda é um desafio a ser enfrentado. Diante da ausência de regulamentação sobre a implementação da tecnologia de inteligência artificial, a presente dissertação tem o objetivo de analisar a necessidade de se criar parâmetros éticos para uma IA de Confiança pautados nos Direitos Humanos, para assim, desenvolver uma regulamentação mais eficaz. Assim, o problema de pesquisa que buscou-se responder é sobre quais os parâmetros que a IA de Confiança deve adotar para o respeito aos Direitos Humanos e os Direitos Fundamentais?

Sendo a hipótese apresentada para o presente problema a eficácia dos parâmetros éticos, a partir de uma interpretação com base na concepção dos Direitos Humanos de Herrera Flores, para a construção da regulamentação de uma inteligência artificial de confiança.

O método utilizado para a abordagem será o hipotético-dedutivo, sendo a hipótese a eficácia dos parâmetros apresentados pelas Orientações Éticas para uma IA de Confiança da Comissão Europeia visando a regulamentação da inteligência artificial. Quanto a sua abordagem, adotou-se o tipo qualitativo pelo entendimento que há relação íntima, porém pouco explorada, entre Judiciário e Inteligência Artificial. No que se refere a técnica de pesquisa, utilizou-se a pesquisa bibliográfica, legislação artigos científicos em meios eletrônicos e notícias em sites oficiais.

Para o desenvolvimento da pesquisa o trabalho foi dividido em três capítulos:

No primeiro capítulo, se abordará brevemente sobre a evolução histórica da IA na sociedade, bem como se apresentará conceitos técnicos referentes a inteligência artificial, como a definição do que é Big Data e algoritmo. Ainda, se apresentará os diferentes tipos de inteligência artificial e as técnicas de aprendizagem, com a finalidade de identificar os tipos de aprendizagem que vem sendo utilizado na sociedade e assim, poder reconhecer as possíveis problemáticas que esta tecnologia vem a trazer, como por exemplo, a obscuridade da *black box*.

O segundo capítulo terá a finalidade de apresentar alguns dos sistemas de inteligência artificial que já foram implementados no sistema judiciário, tanto em âmbito mundial, como no sistema judicial brasileiro. Bem como, analisar as propostas de regulamentação sobre inteligência artificial no mundo e no Brasil.

Por fim, no terceiro capítulo se analisará os parâmetros para uma IA de confiança com enfoque nos Direitos Humanos e na Ética. Dessa forma, pretende-se analisar a partir do conceito da Teoria Crítica dos Direitos Humanos, desenvolvida por Joaquin Herrera Flores, a qual busca contradizer a ideia de universalidade dos Direitos Humanos, propondo que se reconheça a complexidade e pluralidade existe na sociedade.

No que se refere a ética, essa é um fator essencial para a construção de uma IA de confiança, a qual se revela imperiosa a discussão sobre a relação existente entre direito e moral. Tendo em vista, que o direito, como apreciação do fato ou fenômeno social, está subordinado à moral, assim como, à imposição de parâmetros nem sempre perfeitamente racionalizados, os quais se revelam vinculados a pressupostos de senso comum, míticos, religiosos, políticos, etc.

Por fim, se buscará analisar os parâmetros dispostos nas Orientações Éticas para uma IA de Confiança apresentados pela Comissão Europeia, e a partir deste documento construir parâmetros para a regulamentação brasileira.

Além disso, ao abordar qualquer assunto relacionado a IA, não se pode deixar de salientar a necessidade de uma estratégia de governança transnacional, pois seu reflexo ultrapassa as fronteiras estabelecidas pelos Estados Nações, e a partir destes parâmetros construir as regulamentações internas observando as particularidades de cada país.

Sendo assim, é importante salientar, que o que se pretende é uma profunda reflexão acerca das implicações éticas, legais e sociais que a implementação das tecnologias de inteligência artificiais causam na sociedade de forma abrangente a partir da sua implementação no judiciário.

Pois, não restam dúvidas de que as máquinas tendem a consolidação de um sistema único digital e de dados que visa uma maior facilidade e celeridade em um âmbito jurídico. Com isso, para ser possível a instalação de sistemas de inteligência artificial no poder judiciário brasileiro, assim como no âmbito privado, se faz essencial que haja uma confiança entre aqueles que estarão sujeitos a máquina.

1 COMPREENDENDO O APRENDIZADO DE MÁQUINA COM USO DA INTELIGÊNCIA ARTIFICIAL

Este capítulo tem a finalidade de apresentar o uso da inteligência artificial ao longo da história, bem como explicar alguns conceitos essenciais para a compreensão do funcionamento do aprendizado de máquina.

1.1 BREVE EVOLUÇÃO HISTÓRICA DA INTELIGÊNCIA ARTIFICIAL E ALGUNS CONCEITOS RELACIONADOS AO APRENDIZADO DE MÁQUINA

O advento da modernidade teve como um de seus principais avanços a construção de maquinários capazes de substituir, o fazer humano. A produção em massa, marcada pelo atributo do mercado, numa acentuada busca por rapidez, eficiência e lucro, foi produto da Inteligência Artificial na fabricação de diversos produtos e serviços. Aplicações concretas dessa inteligência se tornam cada vez mais frequentes no dia a dia das pessoas e, conseqüentemente, traz à tona, a discussão sobre seu impacto para a personalidade e autonomia pessoal, sendo necessário uma atenção do Direito bem como da comunidade científica para regularizar a mencionada situação.

Este capítulo, portanto, possui o objetivo de permitir uma visão geral da Inteligência Artificial, iniciando com um breve histórico, marcando todos os eventos que originaram essa tecnologia, bem como apresentar a conceituação da Inteligência Artificial e seu impacto na sociedade.

1.1.1 Breve evolução histórica da inteligência artificial

A inteligência artificial vem se tornando cada vez mais presente no cotidiano da sociedade atual, tendo em vista que dentro da rotina social houve um aumento considerável na utilização de assistentes virtuais pelas empresas, os equipamentos eletrônicos possuem reconhecimentos digitais e faciais, e muitos ativam suas funções a partir de comando de voz, até mesmo os anúncios gerados no ambiente virtual são construídos com base no comportamento de acesso de cada pessoa.

Em que pese a contemporaneidade do assunto Inteligência Artificial, as ideias que fomentaram esse tema remontam a tempos muito mais distantes, sendo possível

perceber esse tema nas discussões filosóficas que aconteciam em séculos passados. René Descartes, por exemplo, de acordo com McCorduck¹, embora não tenha sido o primeiro a pensar sobre inteligência humana através de mecanismos, chegou a dividir os atos humanos em atos mecânicos e racionais, apontando com isso as diferenças dos seres humanos para os animais.

Mas segundo Norvig e Russell², estudiosos do tema da Inteligência Artificial, apesar dos filósofos terem demarcado a maioria das ideias sobre essa área, a condição para tornar esse campo como ciência de fato exigiu nível de formalização matemática em três setores fundamentais: a lógica, a computação e a probabilidade.

Foi a partir do século XX que a Inteligência Artificial se manifestou de verdade como ciência, a partir do trabalho realizado por Warren McCulloch e Walter Pitts (1947), na década de 40. Esses dois pesquisadores foram os primeiros a sugerir que redes definidas adequadamente seriam capazes de aprender, visto que eles elaboraram “um modelo de neurônios artificiais, em que cada neurônio se caracterizaria como “ligado” ou “desligado”.³

Porém, foi no contexto do fim da Segunda Guerra Mundial, no ano de 1956, que a nomenclatura Inteligência Artificial foi realmente criada. Conforme Norvig e Russell⁴, a contribuição mais influente para a evolução dessa área foi o desenvolvimento do Colossus, o primeiro computador eletromecânico realizado pela equipe de Alan Turing.

Além disso, Turing publicou o artigo *Computing machinery and intelligence*, em 1950. No mencionado artigo, Alan Turing, elaborou o intitulado “Teste de Turing”, que foi baseado no “Jogo da Imitação”. A ideia por trás do jogo proposto por Turing era verificar o quão próximas as respostas são dadas por um ser humano em comparação a uma máquina. Isto é, ele buscou retratar a capacidade de uma máquina em exibir comportamento inteligente equivalente a um ser humano, ou indistinguível

¹ MCCORDUCK, Pamela. **Machines who think: a personal inquiry into the history and prospects of artificial intelligence**. 2. ed., Massachusetts: A K Peters, 2004, p. 03.

² RUSSELL, Stuart; NORVIG, Peter. **Inteligência Artificial**. 3. ed., Rio de Janeiro: Editora Campus-Elsevier, 2013, p. 09-11.

³ RUSSELL, Stuart; NORVIG, Peter. **Inteligência Artificial**. 3. ed., Rio de Janeiro: Editora Campus-Elsevier, 2013, p. 18.

⁴ RUSSELL, Stuart; NORVIG, Peter. **Inteligência Artificial**. 3. ed., Rio de Janeiro: Editora Campus-Elsevier, 2013, p. 25.

deste.⁵ Conforme Azeredo⁶:

Assim, segundo o Jogo da Imitação de Turing, um humano deve interrogar um computador por via de teletipo – o que hoje seria denominado de mensagens instantâneas – e caso o humano não seja capaz de identificar se está interrogando outro humano ou computador, o computador passa no teste.

Outro marco importante para a evolução da Inteligência Artificial, conforme Russell e Norvig⁷, foi a realização do programa intitulado Logic Theorist, por Allen Newell e Herbert Simon, que foi um programa que utilizava as regras da Lógica para tentar comprovar argumentos. E, em 1956, conforme supracitado, o termo Inteligência Artificial foi cunhado, quando John McCarthy organizou o Dartmouth Summer Research Project on Artificial Intelligence. Um estudo que objetivava sustentar a hipótese inicial de que quaisquer aspectos da inteligência humana podem, a princípio, serem representados por uma máquina.⁸

Ou seja, o que se conhece nos dias de hoje como Inteligência Artificial foi detalhado - nesse Congresso - pelo professor John McCarthy, da Universidade de Stanford que começou a usar o termo na conferência que fez na Faculdade de Dartmouth, em New Hampshire. Em 1961, conforme Urwin⁹, James R. Slage, um norte-americano, redigiu um programa chamado SAINT, capaz de decompor um problema de cálculo em partes menores que seriam facilmente resolvidas por essa máquina.

Já em 1964, Danny Bobrow demonstrou que um computador poderia entender a língua natural. Enquanto em 1965, o alemão Joseph Weizenbaum criou o programa ELIZA, que permitia conversar com usuários. Esse programa, portanto, foi o primeiro software para simulação de diálogos, os chamados "*chatbots*" ou "robôs de

⁵ LIMA, Welton Dias de. Computadores e Inteligência: Uma explicação elucidativa sobre o Teste de Turing. **In: Revista Outras Palavras**, vol. 13, n. 1, 2017. Disponível em: <http://revista.faculdadeprojecao.edu.br/index.php/Projecao5/article/download/756/730>. Acesso em: 24 jul. 2021, p. 04.

⁶ AZEREDO, João Fábio Azevedo e. **Reflexos do emprego de sistemas de inteligência artificial nos contratos**. 2014. 221 fls. Dissertação (Mestrado em Direito) – Universidade de São Paulo, São Paulo, 2014. Disponível em: https://www.teses.usp.br/teses/disponiveis/2/2131/tde-12122014-150346/publico/Dissertacao_reflexos_inteligencia_artificial_contratos_reduzida.pdf. Acesso em: 24 jul. 2021, p. 16.

⁷ RUSSELL, Stuart; NORVIG, Peter. **Inteligência Artificial**. 3. ed., Rio de Janeiro: Editora Campus-Elsevier, 2013, p. 41.

⁸ MONARD, Maria Carolina; BARANAUSKAS, José Augusto. Aplicações de Inteligência Artificial: Uma Visão Geral. **In: Anais – Congresso de Lógica Aplicada à Tecnologia**, São Paulo: Faculdade SENAC de Ciências Exatas e Tecnologia, 2000, p. 341.

⁹ URWIN, Richard. **Artificial Intelligence: The Quest for the Ultimate Thinking Machine**. London: Arcturus, 2016.

conversação", existentes no mundo.¹⁰

Newell e Simon, após realizarem o *Logic Theorist*, prosseguiram com o desenvolvimento do *General Problem Solver*, o prestigiado GPS. Contudo, na década de 70, ocorreu o que os pesquisadores apontam de “inverno da inteligência artificial” visto que nessa época ocorreram poucas novidades para esse campo de estudo somado com diversos cortes de investimentos, que interferiram negativamente na evolução dessa ciência.¹¹

Nos anos 80, porém, sobreveio o lançamento do projeto japonês *Fifth Generation*, que fez com que a pesquisa em Inteligência Artificial voltasse a receber investimentos dos governos e dos órgãos de fomento após um período de desinteresse. Em 1997, ocorreu o que muitos apontam como principal marco histórico da Inteligência Artificial visto que um computador, o *Deep Blue*, da *International Business Machines Corporation* (IBM)¹², venceu um ser humano que no caso era Garry Kasparov, o melhor jogador de xadrez da época, de acordo com a Federação Internacional de Xadrez (FIDE)^{13, 14}.

Deste ponto em diante, a Inteligência Artificial se consolidava como uma indústria, o que pode ser verificado até os dias atuais, movimentando um enorme mercado que não para de se desenvolver a cada dia. Verifica-se, portanto, que foram muitos os avanços dessa área do conhecimento que se deram, principalmente, em razão do progresso da computação, que possibilitou o desenvolvimento dos diversos ramos da Inteligência Artificial.

A seguir, se buscará definir a IA e apresentar as suas principais subáreas.

1.1.2 Conceito de inteligência artificial

Os relatos supracitados descrevem conquistas obtidas com o auxílio de

¹⁰ URWIN, Richard. **Artificial Intelligence: The Quest for the Ultimate Thinking Machine**. London: Arcturus, 2016.

¹¹ RUSSELL, Stuart; NORVIG, Peter. **Inteligência Artificial**. 3. ed., Rio de Janeiro: Editora Campus-Elsevier, 2013, p. 19.

¹² É a maior empresa da área de Tecnologia da Informação no mundo. É uma empresa dos EUA que fabrica e vende hardware e software, oferece serviços de infraestrutura, serviços de hospedagem e serviços de consultoria nas áreas que vão desde computadores de grande porte até a nanotecnologia.

¹³ É uma organização internacional, com sede na Suíça, que conecta as várias federações nacionais de xadrez e atua como órgão dirigente das competições internacionais do esporte.

¹⁴ MONARD, Maria Carolina; BARANAUSKAS, José Augusto. Aplicações de Inteligência Artificial: Uma Visão Geral. *In: Anais – Congresso de Lógica Aplicada à Tecnologia*, São Paulo: Faculdade SENAC de Ciências Exatas e Tecnologia, 2000, p. 345.

ferramentas da Inteligência Artificial, que ao se tornar cada vez mais avançada, fica mais apta a auxiliar os seres humanos em executar tarefas mais complexas de forma mais independente.

Nesse contexto, diante do grande potencial que esse tipo de tecnologia tem de transformar e revolucionar os mais diversos setores da sociedade, começam a surgir questionamentos sobre a necessidade de regulação específica pelo Direito, dos sistemas de IA, de forma a evitar potenciais conflitos e minimizar possíveis danos decorrentes da interação entre as pessoas e esses sistemas. Por isso, para que se possa alinhar normas sociais, éticas e morais à estrutura tecnológica, é preciso, inicialmente, aprender sobre o que é Inteligência Artificial.

Sendo assim, no que tange a IA, não há de se falar em uma unanimidade na conceituação e definição do assunto, todavia, alguns doutrinadores arriscam em dizer que esta falta de delimitação foi o que corroborou com o desenvolvimento deste campo de estudos.

De acordo com Russell e Norvig¹⁵, a IA deveria ser enquadrada em quatro categorias, que se definiram ao longo do tempo, quais sejam: “i) sistemas que pensam como humanos; ii) sistemas que agem como humanos; iii) sistemas que pensam racionalmente e iv) sistemas que agem racionalmente”.

No geral, as linhas de pensamento I e III referem-se ao processo de pensamento e raciocínio, enquanto as II e IV ao comportamento. Além disso, as linhas de pensamento I e II medem o sucesso em termos de fidelidade ao desempenho humano, enquanto na III e IV medem o sucesso comparando-o a um conceito ideal que de inteligência, que se chamará de racionalidade. Um sistema é racional se “faz tudo certo”, com os dados que tem.¹⁶

Cada uma das formas de se enxergar a IA conduziu ao desenvolvimento de escolas de pensamento distintas, cada uma utilizando técnicas e abordagens condizentes com o objetivo específico almejado. A abordagem do teste de Turing, conforme já mencionado, por exemplo, envolve a compreensão da IA como uma máquina que age como um ser humano. Nesse caso, testa-se a capacidade de um *software* se passar por um ser humano a partir de um interrogatório e, se ao final do teste, o interrogador não conseguir identificar quais são as respostas provenientes do

¹⁵ RUSSELL, Stuart; NORVIG, Peter. **Artificial intelligence: a modern approach**. New Jersey: Editora Pearson Prentice Hall, 2010, p. 05.

¹⁶ RUSSELL, Stuart; NORVIG, Peter. **Artificial intelligence: a modern approach**. New Jersey: Editora Pearson Prentice Hall, 2010.

homem e da máquina, será classificado como IA.¹⁷

Outro ensinamento de Norvig e Russell¹⁸, ao classificar IA é de que essa inteligência não se confunde com os sistemas de automação. Em um sistema de automação, segundo os estudiosos, o programador dá à máquina todas as condutas possíveis a serem tomadas, enquanto a Inteligência Artificial possui capacidade de aprender por si própria. A Inteligência Artificial é, nesse aspecto, um mecanismo de acúmulo e representação de conhecimento, que se expande na medida em que coleta mais dados.

Segundo Hartmann Peixoto¹⁹, “a IA pode ser observada como uma estratégia de performance, ou como delegação de funções roboticamente praticáveis, isto é, que envolvam repetição, padrões e volumes em atividades não supervisionadas, mas sempre com fundo ético e responsável”. Brundage caracteriza a IA como um corpo de pesquisa e engenharia utilizando a tecnologia digital a fim de criar sistemas aptos a desempenharem atividades comumente desempenhadas pela inteligência humana. Dando continuidade ao raciocínio, o autor assevera que os efeitos da IA vão além da realização de atividades específicas, mas sim o fato de reunir a tecnologia com propriedades que são originalmente humanas, a exemplo da competência²⁰.

John McCarthy²¹ ao utilizar a expressão Inteligência Artificial, entendia que IA era:

[...] a ciência e a engenharia de se fazer máquinas inteligentes, especialmente programas de computadores inteligentes. Está relacionada à tarefa similar de usar computadores para entender inteligência humana, entretanto IA não necessita estar restrita a métodos que são biologicamente observáveis.²²

Portanto, para Russell e Norvig²³ a inteligência artificial é um campo de estudos da computação que se dedica ao estudo e desenvolvimento de sistemas que possam reproduzir comportamentos inteligentes e cumprir tarefas complicadas com

¹⁷ RUSSELL, Stuart; NORVIG, Peter. **Artificial intelligence: a modern approach**. New Jersey: Editora Pearson Prentice Hall, 2010.

¹⁸ RUSSELL, Stuart; NORVIG, Peter. **Artificial intelligence: a modern approach**. New Jersey: Editora Pearson Prentice Hall, 2010.

¹⁹ PEIXOTO, Fabiano Hartmann; SILVA, Roberta Zumblick Martins. **Inteligência Artificial e Direito**. Curitiba: Alteridade Editora, 2019, p. 19.

²⁰ BRUNDAGE, 2018 apud PEIXOTO, Fabiano Hartmann; SILVA, Roberta Zumblick Martins. **Inteligência Artificial e Direito**. Curitiba: Alteridade Editora, 2019, p. 21.

²¹ MCCARTHY, John. **What is Artificial Intelligence?** Califórnia, Stanford University, 2007. Disponível em: <http://jmc.stanford.edu/articles/whatisai/whatisai.pdf>. Acesso em: 26 jul. 2021, p. 02.

²² It is the science and engineering of making intelligent machines, especially intelligent computer programs. It is related to the similar task of using computers to understand human intelligence, but AI does not have to confine itself to methods that are biologically observable. (texto original)

²³ RUSSELL, Stuart; NORVIG, Peter. **Artificial intelligence: a modern approach**. New Jersey: Editora Pearson Prentice Hall, 2010.

um grau de confiabilidade que é equivalente ou superior ao de um humano.

Patrick Henry Winston²⁴, por sua vez, aduz que existem várias formas de definir a Inteligência Artificial, determinando-a como o estudo da computação que a possibilita de perceber, racionar e agir. Nesse passo, para Urwin²⁵, a IA é:

[...] uma ferramenta construída para ajudar ou substituir o pensamento humano. É um programa de computador, que pode estar numa base de dados ou num computador pessoal ou embutido num dispositivo como um robô, que mostra sinais externos de que é inteligente — como habilidade de adquirir e aplicar conhecimento e agir com racionalidade neste ambiente.

Com o intuito de delimitar o campo de estudo e facilitar a compreensão dos desafios e impactos a serem expostos, a IA nesta pesquisa, será entendida como um conjunto de ferramentas desenvolvidas para que sistemas computacionais possam executar tarefas que requeiram a capacidade racional do ser humano. Ou seja, será adotado o conceito de Inteligência Artificial proposto por Urwin²⁶, apresentado acima, em razão da sua versatilidade.

Diante das diversas tentativas de conceituar a IA, é nítido o que há em comum entre eles, qual seja a criação de tecnologia capaz de equiparar a máquina ao ser humano inserindo, àquela, capacidades, até então, inatas ao homem. Raciocínio, emoção, tomada de decisão, aprendizagem, planejamento são exemplos de aptidões originariamente humanas, onde aplicá-las às máquinas seria um grande avanço tecnológico e revolucionário para a humanidade; no entanto, é importante lembrar que tal realidade ainda é uma meta a ser alcançada.

Por fim, pode-se dizer que as máquinas de Inteligência Artificial desempenham diversas tarefas originalmente advindas da atividade humana, muitas delas executadas de forma mais ágil e prática pelo sistema, do que quando realizadas pelo próprio indivíduo. De forma resumida um sistema inteligente é aquele que apresenta condições humanas que por sua vez demonstra capacidade de raciocinar, planejar, resolver problemas, abstrair ideias, armazenar conhecimento, comunicar-se através de uma linguagem e por fim o mais importante “aprender”.

1.1.3 Big Data

²⁴ WINSTON, Patrick Henry. **Artificial Intelligence**. 3. ed., Boston: Addison-Wesley Publishing Company, 1993, p. 05.

²⁵ URWIN, Richard. **Artificial Intelligence: The Quest for the Ultimate Thinking Machine**. London: Arcturus, 2016, p. 92.

²⁶ URWIN, Richard. **Artificial Intelligence: The Quest for the Ultimate Thinking Machine**. London: Arcturus, 2016.

Para a popularização do conceito, o termo Big Data vem sendo utilizado no mundo todo para veicular a capacidade de analisar e guardar dados em grandes volumes. O Big Data é um conjunto de tecnologias que permite às grandes empresas analisar dados em grande quantidade e traçar uma decisão mais fundamentada a fim de atingir seus objetivos. Os maiores geradores de dados para o Big Data são os seres humanos, pois, estes dados coletados estão presentes em todas as atividades que as pessoas exercem diariamente. Assim, por estar presente em diversas áreas, o conceito de Big Data pode sofrer alterações devido às particularidades de cada campo que é explorado pela humanidade.²⁷

Partindo desse pressuposto, os autores, Guilherme Ataíde Dias e Américo Augusto Nogueira Vieira²⁸, conceituam Big Data através dos conhecimentos de Dumbill, o qual entende que o Big Data é uma enorme quantidade de dados que estão em constante transformação, se mexendo de forma muito rápida, ultrapassando a amplitude dos sistemas de banco de dados convencionais.

Além disso, como demonstra a autora Giovanna Carloni²⁹, assim como o Escritório Executivo do Presidente Barack Obama, ao apresentar um relatório intitulado “Big Data: Seizing Opportunities, Preserving Values”, traz definições sobre o termo Big Data:

Existem muitas definições de “Big Data” que podem diferir dependendo se você é um cientista da computação, um analista financeiro, ou um empreendedor lançando uma ideia para um capitalista de risco. A maioria das definições refletem a grande capacidade da tecnologia de capturar, agregar e processar um sempre crescente volume, velocidade e variedade de dados. Em outras palavras, ‘os dados estão agora disponíveis de forma mais rápida, têm maior cobertura e alcance, e incluem novos tipos de observações e medidas que anteriormente não estavam disponíveis’. Mais precisamente, bases de dados relacionadas ao Big Data são ‘grandes, diversas, complexas, longitudinais e/ou referentes a um conjunto de dados distribuídos e gerados a partir de instrumentos, sensores, transações de internet, e-mail, vídeos, rastreamento do usuário de internet [clickstreams] e/ou todas as outras fontes digitais disponíveis hoje e no futuro.

Ainda se tem a definição de McKinsey Global Institute, apresentada por

²⁷ PEREIRA, Jorge Luís. **Análise Preditiva em Sistemas de Informação no contexto do Big Data**. 2014. 72 fls. Trabalho de Conclusão de Curso (Curso de Sistemas de Informação) – Fundação de Ensino “Eurípes Soares da Rocha”, Marília, 2014.

²⁸ DIAS, Guilherme Ataíde; VIEIRA, Américo Augusto Nogueira. Big Data: questões éticas e legais emergentes. **In: Ciências da Informação**, Brasília, vol. 42, n. 2, p. 174-184, 2013. Disponível em: <http://revista.ibict.br/ciinf/article/view/1380/1558>. Acesso em: 24 jul. 2021.

²⁹ CARLONI, Giovanna Louise Bodin de Saint-Ange Comnène. **Privacidade e Inovação na Era do Big Data**. 2013, 66 fls. Trabalho de Conclusão de Curso (Graduação em Direito) – Fundação Getúlio Vargas, Rio de Janeiro, 2013, p. 07.

Andréia Santos³⁰, que define Big Data como:

É um termo utilizado para descrever um grande volume de dados, em grande velocidade e grande variedade, que requer novas tecnologias e técnicas para capturar, armazenar e analisar seu conteúdo. É utilizado para abrilhantar a tomada de decisão, fornecendo introspecção e descobertas, e suportando e otimizando processos.

Partindo destes conceitos, Andreia Santos³¹, em seu artigo “O impacto do Big Data nas campanhas eleitorais”, aduz sobre três características relevantes sobre o Big Data, denominados de “Vs”:

Conhecidas como os “3 V’s”: (i) volume – a sociedade atual é altamente conectada e tecnológica, todos os dias milhões de transações e comunicações são realizadas online, seja troca de e-mails, mensagens por comunicadores instantâneos, fotos, vídeos, digitalização de documentos, cadastros. Tudo isso corresponde a dados; (ii) velocidade – esses dados são criados de forma acelerada e praticamente instantânea, portanto, atualizadas; e, (iii) variedade – os dados coletados são aleatórios, variados e advêm das mais diversas ferramentas – mídias sociais, celular, gps, sistemas integrados, etc.

Nesse sentido, para Rodrigo Gomes³², há a necessidade da inclusão de mais um “v” nas características que dispõe o Big Data, sendo essa inclusão em razão do termo “veracidade”, pois, neste grande conjunto de dados, podem ocorrer dados imprecisos e até equivocados, necessitando de uma eminente cautela e tecnologia de ponta para não se deixar enganar.

Além dos 4 V’s do Big Data, para alguns autores existe mais uma dimensão que é o valor. Segundo Danielle Parra Ferreira³³: “valor: refere-se à combinação de volume, velocidade, variedade e veracidade, de modo a caracterizar que uma solução Big Data deva representar valor para seu investimento”.

Ainda, sobre a inclusão de mais um V nas características inerentes ao Big Data, discorre sobre o valor, o autor Jorge Luis Pereira³⁴: “o último V que torna a tecnologia Big Data relevante é o V de valor. É importante que todo esse volume,

³⁰ SANTOS, Andréia. **O impacto do Big Data e dos Algoritmos nas Campanhas Eleitorais**. 2017. Disponível em: <https://itsrio.org/wp-content/uploads/2017/03/Andreia-Santos-V-revisado.pdf>. Acesso em: 24 jul. 2021, p. 11.

³¹ SANTOS, Andréia. **O impacto do Big Data e dos Algoritmos nas Campanhas Eleitorais**. 2017. Disponível em: <https://itsrio.org/wp-content/uploads/2017/03/Andreia-Santos-V-revisado.pdf>. Acesso em: 24 jul. 2021, p. 11.

³² GOMES, Rodrigo Dias de Pinho. **Desafios à Privacidade: Big data, consentimento, legítimos interesses e novas formas de legitimar o tratamento de dados pessoais**. 2017. Disponível em: <https://itsrio.org/wp-content/uploads/2017/03/Rodrigo-Gomes.doc-B.pdf>. Acesso em: 24 jul. 2021.

³³ FERREIRA, Danielle Parra. **Análise Comparativa do uso do Big data em redes sociais: Um estudo sobre o posicionamento de juízes e advogados**. 2016. 126 fls. Trabalho de Curso (Curso de Gestão da Informação) – Universidade Federal do Paraná, Curitiba, 2016, p. 21.

³⁴ PEREIRA, Jorge Luís. **Análise Preditiva em Sistemas de Informação no contexto do Big Data**. 2014. 72 fls. Trabalho de Conclusão de Curso (Curso de Sistemas de Informação) – Fundação de Ensino “Eurípes Soares da Rocha”, Marília, 2014, p. 22.

variedade, velocidade e veracidade possam gerar valor”.

Ao falar sobre os dados a partir de onde eles são coletados, como por exemplo, nas empresas, também existem os dados obtidos através de toda a rede (internet). Entre eles encontram-se os meios de comunicação sociais, que são explorados pelos usuários de smartphones e outros aparelhos eletrônicos. Adendo que faz o autor César Taurion³⁵:

Integram o chamado Big Data o conteúdo de 640 milhões de sites, dados de seis bilhões de celulares e os três bilhões de comentários feitos diariamente no Facebook. Variedade porque estamos tratando tanto de dados textuais estruturados quanto não estruturados como fotos, vídeos, e-mails e tuítes. E velocidade, porque muitas vezes precisamos responder aos eventos quase que em tempo real. Ou seja, estamos falando de criação e tratamento de dados em volumes massivos. A questão do valor é importante. Big Data só faz sentido se o valor da análise de dados compensar o custo de sua coleta, armazenamento e processamento [...].

Como visto, o termo Big Data engloba diversos conceitos. As características do Big Data são extremamente amplas e cada autor enxerga uma diferente maneira de conceituá-las, o que acaba por agregar-lhe ainda mais valor.

1.1.4 Algoritmo

Algoritmo, de maneira simplificada, é a descrição de um conjunto de ações que, obedecidas, resultam numa sucessão finita de passos atingindo um objetivo. Sendo assim, “[...] uma sequência de raciocínios, instruções ou operações para alcançar um objetivo, sendo necessário que os passos sejam finitos e operados sistematicamente”.³⁶

Essa sequência de passos deve ser finita. Um exemplo clássico é uma receita de bolo, onde há um input (os ingredientes), após, ocorre o processamento, realizado com uma sequência de passos bem definidas, não ambíguas e finitas (modo de preparo), para por fim, obter-se o output (o bolo pronto). A diferença é que, quando se trata de um algoritmo de inteligência artificial, estes são muito mais complexos, além de serem capazes de tomar decisões sozinhos.

Conforme afirma Wang³⁷, o conceito de algoritmo, por si só, não é

³⁵ TAURION, César. **Big data**. Rio de Janeiro: Brasport Livros e Multimídia Ltda, 2013, p. 32.

³⁶ ROCKCONTENT. **Saiba como funciona um algoritmo e conheça os principais exemplos existentes no mercado**. 2019. Disponível em: <https://rockcontent.com/br/blog/algoritmo/>. Acesso em: 05 set. 2021.

³⁷ Texto original: “The intuitive notion of an algorithm is quite vague. For example, what is a rule? We would like the rules to be mechanically interpretable, that is, so that a machine can understand the rule

suficientemente preciso:

A noção intuitiva de um algoritmo é bastante vaga. Por exemplo, o que é uma regra? Gostaríamos que as regras fossem mecanicamente interpretáveis, ou seja, de modo que uma máquina possa entender a regra (instrução) e realizá-la. Em outras palavras, precisamos especificar uma linguagem para descrever algoritmos que seja abrangente o suficiente para descrever todos os procedimentos mecânicos e ainda simples o bastante para ser interpretado por uma máquina. (Tradução nossa)

Segundo Rômulo Soares Valentini³⁸, inicialmente, é necessário estabelecer o mecanismo de entrada de dados (input). Um algoritmo deve ter um ou mais meios para recepção dos dados a serem analisados. Em uma máquina computacional, a informação deve ser passada para o computador em meio digital (bits). Do mesmo modo, é necessário ter um mecanismo para a saída ou retorno dos dados trabalhados (output). Um algoritmo deve ter um ou mais meios para retorno dos dados, os quais devem estar relacionados de modo específico com o input. Por exemplo, um algoritmo de uma calculadora que receba as informações para somar 2+2 (input) irá retornar como resultado o número 4 (output). O output decorre do input, sendo papel do algoritmo fornecer o retorno dos dados corretos a partir dos dados de entrada. Uma vez que o algoritmo não faz nenhum juízo de valor para além de sua programação, é necessário que a relação de “correção” entre o input e o output seja definida de modo preciso e sem ambiguidade. Por isso, os algoritmos precisam ter cada passo de suas operações cuidadosamente definido. Assim, cada passo da tarefa computacional deve seguir um roteiro de tarefas pré-determinado e o programa (computação dos dados) deve terminar depois que o roteiro seja cumprido. O algoritmo tem que ser finito, ou seja, entregar algum retorno (output) após cumpridos todos os passos estabelecidos. Para cumprir a tarefa adequadamente, cada operação que o algoritmo tiver que realizar deve ser simples o suficiente para que possa ser realizada de modo exato e em um tempo razoável (finito) por um ser humano usando papel e caneta. Conclui-se, desse modo, que um o algoritmo é um plano de ação pré-definido a ser seguido pelo computador, de maneira que a realização contínua de pequenas tarefas simples possibilitará a realização da tarefa solicitada sem novo dispêndio de trabalho humano.

(instruction) and carry it out. In other words, we need to specify a language for describing algorithms that is comprehensive enough to describe all mechanical procedures and yet simple enough to be interpreted by a machine”. WANG, Hao. **From Mathematics to Philosophy**. London: Routledge, Kegan & Paul, 1974, p. 91.

³⁸ VALENTINI, Romulo Soares. **Julgamento por computadores?** As novas possibilidades da juscibernética no século XXI e suas implicações para o futuro do direito e do trabalho dos juristas. 2017. 152 fls. Tese (Doutorado em direito) – Universidade Federal de Minas Gerais, Belo Horizonte, 2017, p. 42.

Portanto, é possível concluir que a construção do algoritmo é abstrata, até que ele seja codificado em um programa por meio de uma linguagem de programação, onde se pode testá-lo e concretizá-lo em um processo computacional, em que a partir da sua utilização passa a ter as características de repetibilidade.

Nesse sentido, no momento atual, é possível afirmar que os algoritmos, por meio do *machine learning* e *deep learning*, conseguem induzir premissas, isto é, construir hipóteses com base em um determinado conjunto de dados.³⁹

1.1.5 Black Box

Além disso, há outra característica da Inteligência Artificial que possuem importante reflexo jurídico e as quais merecem ser mencionadas. Uma delas é a chamada *black box* da Inteligência Artificial, que são decisões independentes que muitas vezes serão, ininteligíveis ou opacas, em virtude da estrutura que baseia seu próprio funcionamento. Isto é, muito se fala da *black box* da Inteligência Artificial, como a “caixa preta”, cujo interior não pode ser visualizado, onde ocorre o processamento do sistema.⁴⁰

Em determinadas maneiras de aplicação da IA, as informações externas que são alimentadas ao sistema são direcionadas a uma rede neural artificial que processa os dados e, em seguida, distribui os comandos necessários para operar o sistema no mundo físico ou virtual. No entanto, na maioria das vezes, não é possível, de maneira técnica, refazer o caminho lógico tomado pela rede neural do sistema para saber o porquê de tal operação.⁴¹

Interessante apontar que uma das soluções que vem sendo defendida pela academia para sanar a entrave supracitada, e que é inclusive objeto do Projeto de Lei nº 2018/49 da cidade de Nova Iorque, Estados Unidos, é estabelecer diretrizes para a criação de mecanismos de transparência nesses sistemas, com o intuito de facilitar

³⁹ SHIPPERS, Laurianne-Marie. **Algoritmos que discriminam**: uma análise jurídica da discriminação no âmbito das decisões automatizadas e seus mitigadores. 2018. 57 fls. Monografia (Graduação em Direito) – Escola de Direito de São Paulo, Fundação Getúlio Vargas, São Paulo, 2018. Disponível em: <http://bibliotecadigital.fgv.br/dspace/bitstream/handle/10438/29878/Algoritmos%20que%20discriminam%20-%20Laurianne-Marie%20Schippers.pdf?sequence=1&isAllowed=y>. Acesso em: 05 set. 2021.

⁴⁰ KNIGHT, Will. The Dark Secret at the Heart of AI. **In: MIT Technology Review**, 2017. Disponível em: <https://www.technologyreview.com/2017/04/11/5113/the-dark-secret-at-the-heart-of-ai/>. Acesso em: 30 jul. 2021.

⁴¹ KNIGHT, Will. The Dark Secret at the Heart of AI. **In: MIT Technology Review**, 2017. Disponível em: <https://www.technologyreview.com/2017/04/11/5113/the-dark-secret-at-the-heart-of-ai/>. Acesso em: 30 jul. 2021.

a percepção dessas conexões causais. Ou seja, como criação da IA se dá frequentemente de forma difusa e seu funcionamento ocorre de maneira inexplicável e opaca, dentro de uma *black box* inacessível tanto aos usuários do sistema quanto aos seus desenvolvedores, estudiosos da área vem defendendo a implementação de diretrizes para criação de um sistema mais transparente.⁴²

Dito isso, é possível perceber que os sistemas de IA são inteligentes em um sentido formal, uma vez que tratamos dos meios de funcionamento da máquina e da forma como ela age em determinadas situações, independente da matéria nela implementada. Desta forma, são factíveis de alteração e aperfeiçoamento de seu comportamento a partir de experiências a ela implantada.

1.1.6 Processamento de Linguagem Natural (PLN)

A área de processamento de linguagem natural (PLN) se encontra posicionada dentro da área de Inteligência Artificial, por fazer uso dos mecanismos desse segundo campo de estudo, para realizar a interpretação e a compreensão dos textos. Consiste em um campo interdisciplinar que inclui a IA, a ciência cognitiva, o processamento de informações e a linguística e busca habilitar as máquinas a processar línguas humanas de forma inteligente, tornando a IA mais natural de forma a facilitar a “conversação” e a solução do problema fazendo-se mais fluída e descomplicada.

De acordo com Guilherme Busato Vecchi e Thábata Pontes Lazarou Risso⁴³, o Processamento de Linguagem Natural é um campo da Inteligência Artificial que visa possibilitar que as máquinas possam ler, entender e induzir significados a partir da linguagem natural humana. Atualmente, esse campo de estudo tem ganhado cada vez mais relevância, pois o acesso aos dados e ao poder computacional tem se tornado cada vez mais fácil e barato, possibilitando a criação de soluções na área de saúde, finanças, indústria, etc.

⁴² BIRD, Sarah; BAROCAS, Solon; CRAWFORD, Kate; DIAS, Fernando; WALLACH, Hanna. Exploring or Exploiting? Social and Ethical Implications of Autonomous Experimentation in AI. **In: Workshop on Fairness, Accountability, and Transparency in Machine Learning**. 2016. Disponível em: <https://www.microsoft.com/en-us/research/publication/exploring-or-exploiting-social-and-ethical-implications-of-autonomous-experimentation-in-ai/>. Acesso em: 31 jul. 2021.

⁴³ VECCHI, Guilherme Busato; RISSO, Thábata Pontes Lazarou. **Solução do Desafio de Esquemas de Winograd em Português Brasileiro**. 2020. 94 fls. Monografia (Graduação em Engenharia) – Escola Politécnica da Universidade de São Paulo, São Paulo, 2020.

Conforme lecionam Deng e Yang⁴⁴, o desenvolvimento do campo do processamento de linguagem natural pode ser descrito em três ondas principais: racionalismo, empirismo e *deep learning*. A primeira onda, com abordagem racionalista, defendia que se definissem regras manufaturadas para a incorporação de conhecimento nos sistemas de processamento de linguagem, assumindo que esse conhecimento na mente humana era previamente fixado por herança genética. A segunda onda, de abordagem empírica, presume que a absorção sensorial e a observação dos dados da linguagem são exigidos e suficientes para permitir a mente a aprender a detalhada estrutura da linguagem natural. Como resultado dessa onda, foram desenvolvidos modelos probabilísticos para descobrir regularidades das linguagens a partir de uma coleção de escritos. A terceira onda, do *deep learning*, explora modelos hierárquicos de processamento não linear inspirados pelos sistemas neurais biológicos para aprender representações intrínsecas dos dados de linguagem, de forma a simular as habilidades cognitivas humanas.

O uso de NLP está cada vez mais presente, uma vez que consegue criar uma interface mais natural para a entrada de dados em um computador. Essas entradas mais naturais podem ocorrer de diversas maneiras, indo desde uma pessoa falando e dando ordens para o computador até a interpretação de textos que foram escritos por humanos e não possuem nada além de texto para ser lido, algo que acontece muito comumente na internet atualmente.

Por fim, a partir do conhecimento breve sobre a Inteligência artificial e como ocorre o aprendizado de máquinas, quando tratada em um âmbito forense, a inteligência artificial ainda se encontra em fase embrionária, porém já é aplicada de diversos formatos adotando estes sistemas supervisionados. Programas deste tipo foram implementados em escritório jurídicos visando tornar o trabalho de advogados mais simples, eficientes e organizados. Também estão presentes no poder judiciário em diversos tribunais regionais e superiores no auxílio de tomada de decisões, pesquisa de jurisprudência, automatização e organização de peças processuais. Até o momento esses mecanismos estão sendo usado de forma auxiliar nas decisões dos juristas como uma forma de consultoria, todavia, não há discussão quando o avanço dessa tecnologia no Direito, com isso, é de extrema necessidade o debate acerca do tema e de seus possíveis impactos.

⁴⁴ LI, Deng; LIU, Yang. **Deep learning in natural language processing**. Suíça: Springer, 2018.

1.2 INTELIGÊNCIA ARTIFICIAL FRACA, FORTE E SUPERINTELIGÊNCIA

Para que se tenha uma melhor compreensão de como aplicar a inteligência artificial, não basta apenas reconhecer que uma máquina seja considerada “inteligente”, é fundamental compreender qual o alcance dessa inteligência até para saber qual seu grau de eficiência quando empregada.

Partindo desse princípio, no presente trabalho se levará em consideração três tipos de Inteligência Artificial, classificadas, respectivamente, como fraca, forte e superinteligente.

1.2.1 Inteligência Artificial Fraca

A inteligência artificial fraca, seria a *Artificial Narrow Intelligence* (Inteligência Artificial Estreita), consiste de algoritmos especializados em resolver problemas em uma área e/ou um problema específico. Aqui os sistemas armazenam uma grande quantidade de dados e os algoritmos são capazes de realizar tarefas complexas, porém sempre focadas no objetivo para o qual foram desenvolvidos.⁴⁵

Sendo assim, o conceito de IA fraca refere-se a simples tentativa do sistema em replicar ou duplicar as funções aprendidas e para isso a maioria das tarefas é o suficiente, assim a IA fraca é a simulação da tomada de decisão.

Nesse passo, de acordo com Fernanda Pacheco Amorim⁴⁶:

Por tudo isso que se afirma que a inteligência artificial fraca não tem capacidade de raciocínio próprio, pelo contrário, somente utiliza-se de conhecimentos e insumos repassados pelos seres humanos, dizendo-se, inclusive, que enquanto a inteligência artificial for capaz apenas de escrever um poema organizando palavras e não tem consciência de que escreveu um poema, a partir de sentimentos e emoções, ela continuará sendo incapaz de raciocínio autônomo e consequentemente será considerada uma inteligência artificial fraca.

Corroborando com este entendimento, Searle⁴⁷ explica que a inteligência artificial chamada de fraca “nos dá ferramentas muito potentes”, isto é, “nos permite

⁴⁵ RICHARDS, Neil M.; SMART, William D. **How should the law think about robots?**. 2016. Disponível em: http://robots.law.miami.edu/wp-content/uploads/2012/03/RichardsSmart_HowShouldTheLawThink.pdf. Acesso em: 30 jul. 2021.

⁴⁶ AMORIM, Fernanda Pacheco. **Respeita as Mina: Inteligência artificial e violência contra a mulher**. 1. ed., Florianópolis: EMais Editora, 2019, p. 98.

⁴⁷ SEARLE, John R.. *Minds, Brains and Programs*, **In: The Behavioral and Brain Sciences**, vol. 3, n. 3, p. 417-457, Cambridge: Cambridge University Press, 1980, p. 417.

formular e testar hipóteses de forma mais rigorosa e precisa”, porém, ela depende da inserção de conhecimento fornecido pelo ser humano que a programa, sendo que a máquina não é capaz de produzir raciocínios próprios, autônomos.

Alguns exemplos de IA Fraca a partir de 2018 no Sistema Jurídico são:

- Dra. Luzia Robô Advogada brasileira;
- Assistente Digital do Promotor;
- As Robôs Alice, Sofia e Mônica;
- E, o Robô Victor;⁴⁸

Portanto, na hipótese da IA Fraca, admite-se apenas que uma simulação de consciência seja viável, mas a consciência de fato não será alcançável para uma máquina.

1.2.2 Inteligência Artificial Forte

Já a inteligência artificial forte, é a *Artificial General Intelligence* (Inteligência Artificial Geral), que mimetiza a mente humana e tem várias habilidades, ou seja, demonstra uma capacidade além daquelas relacionadas ao armazenamento de programação, a cognitiva, apresentada pela capacidade de pensar, aprender e tomar decisões de forma autônoma, desenvolvendo tarefas que requeiram alguma inteligência de maneira igual ou melhor que um humano.⁴⁹

Um exemplo de IA Forte é o Robô cognitivo, o Ross Advogado Assistente, que atua em um escritório jurídico americano.⁵⁰

Portanto, conforme Russel e Norvig⁵¹ apresentam, o termo Inteligência Artificial Forte é usado para enfatizar o ambicioso objetivo de se criarem sistemas inteligentes com competências amplas, cuja amplitude de aplicação seria ao menos comparável com as múltiplas tarefas que os humanos podem realizar.

⁴⁸ COSTA, Suzana Rita da. **A Contribuição da Inteligência Artificial na Celeridade dos Trabalhos Repetitivos no Sistema Jurídico**. 2020. 72 fls. Dissertação (Mestrado em Mídia e Tecnologia) – Universidade Estadual Paulista Júlio de Mesquita Filho – Unesp, Bauru, 2020, p. 29.

⁴⁹ RICHARDS, Neil M.; SMART, William D. **How should the law think about robots?**. 2016. Disponível em: http://robots.law.miami.edu/wp-content/uploads/2012/03/RichardsSmart_HowShouldTheLawThink.pdf. Acesso em: 30 jul. 2021.

⁵⁰ COSTA, Suzana Rita da. **A Contribuição da Inteligência Artificial na Celeridade dos Trabalhos Repetitivos no Sistema Jurídico**. 2020. 72 fls. Dissertação (Mestrado em Mídia e Tecnologia) – Universidade Estadual Paulista Júlio de Mesquita Filho – Unesp, Bauru, 2020, p. 30.

⁵¹ RUSSELL, Stuart; NORVIG, Peter. **Inteligência Artificial**. 3. ed., Rio de Janeiro: Editora Campus-Elsevier, 2013.

Segundo Searle⁵², na chamada inteligência artificial “forte”, “o computador não é uma mera ferramenta no estudo da mente, ao contrário, o computador adequadamente preparado é realmente uma mente, no sentido de que os computadores que recebem os programas certos poderiam estar, literalmente, preparados para compreender e ter outros estados cognitivos”. A inteligência “forte”, portanto, seria aquela capaz de criar consciência, simulando raciocínios complexos e emitindo opiniões autônomas, independente da interferência constante do ser humano.

Salienta-se, contudo, que Searle apesar de reconhecer que a Inteligência forte se caracteriza pela capacidade de a máquina criar consciência, o autor defende este conceito de modo crítico. Sendo assim, Searle afirma veementemente que uma Inteligência Artificial forte, não teria exatamente uma consciência própria, pois a IA forte responde perguntas de forma que alguém de fora, humano, sente que ela tem consciência, quando na verdade não existe um processo interno consciente, o cérebro feito de matéria biológica é parte do que gera nossa consciência (humana) e, sem ele uma consciência de verdade não existiria.⁵³

1.2.3 Inteligência Artificial Superinteligente

Por fim, a terceira seria a Inteligência Artificial Superinteligente, *Artificial Super Intelligence* (Inteligência Super Artificial), que pondera que seria mais inteligente até mesmo que o cérebro humano em diversas áreas. Alguns pesquisadores apontam que não seria e nem será possível a Inteligência Artificial alcançar este último patamar haja visto acreditarem ser inviável essa inteligência possuir maior capacidade cognitiva que os seres humanos.⁵⁴

O exemplo mais simples da superinteligência rápida seria uma emulação completa do cérebro executada em um hardware veloz. Uma emulação operando com uma velocidade 10 mil vezes maior que a de um cérebro biológico seria capaz de ler um livro em alguns segundos e escrever uma tese de doutorado em uma tarde. Se a

⁵² SEARLE, John R.. Minds, Brains and Programs, *In: The Behavioral and Brain Sciences*, vol. 3, n. 3, p. 417-457, Cambridge: Cambridge University Press, 1980, p. 417.

⁵³ SEARLE, John R.. Minds, Brains and Programs, *In: The Behavioral and Brain Sciences*, vol. 3, n. 3, p. 417-457, Cambridge: Cambridge University Press, 1980, p. 417.

⁵⁴ RICHARDS, Neil M.; SMART, William D. **How should the law think about robots?**. 2016. Disponível em: http://robots.law.miami.edu/wp-content/uploads/2012/03/RichardsSmart_HowShouldTheLawThink.pdf. Acesso em: 30 jul. 2021.

velocidade fosse 1 milhão de vezes maior, uma emulação poderia realizar o trabalho intelectual de um milênio em apenas um dia de trabalho. Para uma mente rápida como essa, os eventos do mundo externo parecem acontecer em câmera lenta.⁵⁵

Conforme explica Fabiano Hartmann Peixoto e Roberta Zumblick Martins Silva⁵⁶, existe muita preocupação e especulação a respeito das implicações dos computadores serem mais inteligentes do que as pessoas. Há quem preveja que uma Inteligência Artificial suficientemente inteligente possa ser atarefada de desenvolver sistemas ainda melhores, levando a uma “explosão de inteligência” ou a “singularidade” nas quais os computadores seriam muito mais inteligentes do que os humanos. Em uma visão distópica, essas máquinas superinteligentes excederiam as habilidades humanas de compreensão e controle.

Esse discurso apocalíptico se fortaleceu a partir dos dizeres de Nick Bostrom⁵⁷, o qual afirma em tom de alerta que,

[...] se algum dia construirmos cérebros artificiais capazes de superar o cérebro humano em inteligência geral, então essa nova superinteligência poderia se tornar muito poderosa. E, assim como o destino dos gorilas depende mais dos humanos do que dos próprios gorilas, também o destino de nossa espécie dependeria das ações da superinteligência de máquina.

Por outro lado, Fernando A. G. Alcoforado⁵⁸ ressalta que se o ser humano não evoluir também através da tecnologia, ela não terá servido para nada. Da mesma forma que pode surgir a singularidade tecnológica com a superinteligência artificial, poderá surgir, também, a singularidade humana com a formação de superhomens e supermulheres que possam sobreviver a doenças e pandemias, bem como sejam capazes biologicamente de sair do planeta Terra e realizar viagens espaciais. A singularidade humana é alcançada com o transhumanismo que é uma filosofia que tem como objetivo melhorar a condição humana a partir do uso de ciência e da tecnologia (biotecnologia, nanotecnologia e neurotecnologia) para aumentar a capacidade cognitiva e superar limitações físicas e psicológicas dos seres humanos.

⁵⁵ BOSTROM, Nick. **Superinteligência: caminhos, perigos e estratégias para um novo mundo**. Rio de Janeiro: DarkSide Books, 2018, p. 109.

⁵⁶ PEIXOTO, Fabiano Hartmann; SILVA, Roberta Zumblick Martins. **Inteligência Artificial e Direito**. Curitiba: Alteridade Editora, 2019, p. 79-80.

⁵⁷ BOSTROM, Nick. **Superinteligência: caminhos, perigos e estratégias para um novo mundo**. Rio de Janeiro: DarkSide Books, 2018, p. 15.

⁵⁸ ALCOFORADO, Fernando A. G. **Os benefícios e os riscos da singularidade tecnológica baseada na Superinteligência Artificial**. 2020. Disponível em: https://www.portalsaudenoar.com.br/wp-content/uploads/2020/12/OS_BENEFICIOS_E_OS_RISCOS_DA_SINGULARIDA.pdf. Acesso em: 05 set. 2021.

1.3 TÉCNICAS DE APRENDIZAGEM EM INTELIGÊNCIA ARTIFICIAL

O fenômeno da Inteligência Artificial, como visto anteriormente, traz conceitos incomuns à maior parcela da sociedade e, por isso, se faz necessário a elucidação de algumas concepções, como por exemplo, os termos como machine learning, deep learning e redes neurais artificiais, todos relacionados as técnicas de aprendizagem para a construção de uma IA.

Figura 1 – Concepções da Inteligência Artificial



Fonte: Revista Programar⁵⁹

No entanto, é necessário salientar, conforme demonstra figura acima, que um depende do outro para que algoritmos se desenvolvam como inteligência artificial.

1.3.1 Machine learning

O termo *machine learning*, consiste na capacidade que o sistema tem de adquirir conhecimento próprio por meio da extração de padrões a partir de dados brutos - coletados por meio do uso de algoritmos.⁶⁰ Logo, trata-se de um ramo da Inteligência Artificial baseado na ideia de que sistemas podem aprender com dados, identificar padrões e tomar decisões com o mínimo de intervenção humana.

⁵⁹ SARAIVA, Sérgio. Deep Learning passo a passo. *In: Revista Programar*. 2018. Disponível em: <https://www.revista-programar.info/artigos/deep-learning-passo-passo/>. Acesso em: 08 set. 2021.

⁶⁰ MAGRANI, Eduardo. **Entre dados e robôs: ética e privacidade na era da hiperconectividade**. 2. ed., Porto Alegre: Arquipélogo Editorial, 2019.

Nesse sentido, Katti Faceli et al.⁶¹ disserta que a utilização de *machine learning* se dá de modo “que computadores são programados para aprender com a própria experiência passada”. Por isso, foram desenvolvidos algoritmos capazes de obter conclusões a partir de um conjunto de exemplos, assim aprendendo a deduzir uma hipótese ou função capaz de resolver um problema baseado em dados que demonstram iminência do problema a ser solucionado.

Do mesmo modo, Stuart Russell e Peter Norvig⁶² definem *machine learning* como o ramo da IA que estuda formas de fazer com que os computadores melhorem sua performance com base na experiência. Segundo os autores, há uma ideia equivocada de que *machine learning* seja um ramo novo, com o objetivo de substituir a IA em larga escala. Ao contrário, *machine learning* sempre foi um assunto central em IA.

Um dos objetivos desse ramo da IA é possibilitar que os computadores aprendam “sozinhos”. Um algoritmo de machine learning permite que esta identifique padrões nos dados sob exame, construa modelos que expliquem o “mundo” e preveja coisas sem regras e modelos explicitamente pré-programados.⁶³

Explicando de uma maneira lúdica, Ana Lídia Lira Ribeiro⁶⁴ demonstra a forma que funcionam os modelos de *Machine Learning*:

Dessa forma, os modelos de Machine Learning procuram por padrões nos dados com que são alimentados e tiram conclusões, da mesma forma que os seres humanos fazem ao aprender algo novo. Como exemplo, tem-se uma criança que está aprendendo a diferenciar animais. Os pais falam que tal animal é um “au au”. Logo a criança começa a procurar padrões, identificando as características que está vendo: quatro patas, peludo, tem um rabo e associando elas à informação recebida de que aquilo é um “au au”. Depois a criança ao ver um gato pela primeira vez irá, provavelmente, chegar a conclusão após analisar e visualizar os mesmos padrões que aquilo é um “au au”, contudo ao ser corrigida pelos pais que dizem que é um “miau”, mais uma informação entrou no banco de dados mental da criança, ajudando-a assim, a formar outros padrões diferenciando cachorros e gatos.

Portanto, o *machine learning* nasceu do reconhecimento de padrões e da teoria de que computadores podem aprender sem serem programados para realizar

⁶¹ FACELI, K.; LORENA, A. C.; GAMA, J.; CARVALHO, A. C. P. L. F.. **Inteligência Artificial: Uma Abordagem de Aprendizagem de Máquina**. Rio de Janeiro: LTC, 2011.

⁶² RUSSELL, Stuart; NORVIG, Peter. **Inteligência Artificial**. 3. ed., Rio de Janeiro: Editora Campus-Elsevier, 2013.

⁶³ MAINI, Vishal; SABRI, Samer. **Machine Learning for Humans**. 2019. Disponível em: <https://everythingcomputerscience.com/books/Machine%20Learning%20for%20Humans.pdf>. Acesso em: 08 set. 2021.

⁶⁴ RIBEIRO, Ana Lídia Lira. **Discriminação em algoritmos de inteligência artificial: uma análise acerca da Igpd como instrumento normativo mitigador de vieses discriminatórios**. 2021. 61 fls. Monografia (Graduação em Direito) – Universidade Federal do Ceará, Fortaleza, 2021, p. 20.

tarefas específicas. Alguns exemplos bem conhecidos de aplicações de *machine learning*, são os carros autônomos da Uber⁶⁵, as ofertas recomendadas pela Amazon⁶⁶ e as indicações de filme pela Netflix⁶⁷. E, cada vez mais, grandes empresas frente ao mercado têm aumentado o interesse na adoção do *machine learning* como parte de uma estratégia analítica maior, para incorporá-lo na infraestrutura de dados de produção.

Essa tecnologia possibilitou que computadores pudessem lidar com problemas que exigem conhecimento do mundo real e tomar decisões que aparentam subjetividade. Goodfellow, Bengio e Courville citam como exemplo um algoritmo de *machine learning* simples, denominado regressão lógica, que pode determinar casos médicos nos quais se recomenda um parto cesariana pela avaliação de fatores de risco. Outro exemplo de algoritmo de *machine learning* simples é chamado e Naive Bayer, o qual é capaz de separar e-mails legítimos de e-mails de spam.⁶⁸

Todavia, cabe destacar que as máquinas não apresentam isoladamente a capacidade de aprendizado, elas só podem fazer o que seus programadores comandam que façam.

Neste sentido, esclarecem Dierle Nunes e Ana Luíza Marques⁶⁹:

Inicialmente, importante consignar que os mecanismos de inteligência artificial dependem de modelos, os quais consistem em representações abstratas de determinado processo, sendo, em sua própria natureza, simplificações de nosso mundo real e complexo. Ao criar um modelo, os programadores devem selecionar as informações que serão fornecidas ao sistema de IA e que serão utilizadas para rever soluções e/ou resultados futuros.

Ainda acerca da *machine learning*, ressalta-se que alguns autores a dividem em três tipos: supervisionado, não supervisionado e por reforço.

⁶⁵ É uma empresa multinacional norte americana, que presta serviços eletrônicos na área do transporte privado urbano, através de um aplicativo que permite a busca por motoristas que façam sua viagem baseada na localização do cliente.

⁶⁶ É uma empresa de tecnologia norte-americana que se concentra no e-commerce, computação em nuvem, streaming e inteligência artificial.

⁶⁷ É uma provedora global de filmes e séries de televisão via *streaming*.

⁶⁸ PEIXOTO, Fabiano Hartmann; SILVA, Roberta Zumblick Martins. **Inteligência Artificial e Direito**. Curitiba: Alteridade Editora, 2019, p. 89.

⁶⁹ NUNES, Dierle; MARQUES, Ana Luiza Pinto Coelho. Inteligência Artificial e direito processual: vieses algorítmicos e os riscos de atribuição de função decisória às máquinas. **In: Revista de Processo**, São Paulo, vol. 285, p. 421-447, 2018. Disponível em: https://www.academia.edu/37764508/INTELIG%C3%8ANCIA_ARTIFICIAL_E_DIREITO_PROCESSUAL_VIESES_ALGOR%C3%8DTMICOS_E_OS_RISCOS_DE_ATRIBUI%C3%87%C3%83O_DE_FUN%C3%87%C3%83O_DECIS%C3%93RIA_%C3%80S_M%C3%81QUINAS_Artificial_intelligence_and_procedural_law_algorithmic_bias_and_the_risks_of_assignment_of_decision_making_function_to_machines. Acesso em: 30 jul. 2021, p. 04.

1.3.1.1 Aprendizagem supervisionada

Algoritmos de aprendizagem supervisionada funcionam de forma comparativa, costumam ser utilizadas na realização de tarefas de classificação e de regressão. O sistema recebe previamente o algoritmo e seu padrão a ser seguido. Ao receber esses dados, os analisa e então compara sua repetição com o conjunto já fornecido, agrupando-os em seus grupos específicos e chegando à solução esperada. Este método é utilizado, de modo exemplificativo, para a associação e reconhecimento de números escritos e números caligrafados, normalmente utilizados em fotografias de cheques.

Para melhor compreender como funciona o algoritmo de aprendizagem supervisionada, Fabiano Hartmann Peixoto e Roberta Zumblick Martins Silva⁷⁰ explicam que:

[...] nos problemas de aprendizagem supervisionada, começa-se por um conjunto de dados (dataset) de treinamento, no qual os exemplos estão associados com uma rotulagem correta. Maine e Sabrini (2017) dão o exemplo de um algoritmo que deve aprender a reconhecer números escritos em caligrafia humana. Milhares de fotografias de números escritos à mão são expostos ao algoritmo junto ao número correto correspondente. O algoritmo aprenderá a relação entre as imagens e seus números associados, para depois aplicar a relação apreendida a imagens completamente novas e não rotuladas. É por isto que é possível fazer o depósito de um cheque tirando apenas uma foto no celular. Esta é uma tarefa de classificação, de predição de uma classe, na qual o algoritmo consegue reconhecer a que classe (no caso, a qual o número cardinal) a caligrafia corresponde.

Outro exemplo de aprendizagem supervisionada pode ser o histórico de músicas que uma pessoa gosta ou não, a músicas tem de entradas de intensidade da música, gênero da música e tom, e a classe de saída conhecida dessa música é se a pessoa gostou ou não gostou da música. Com isso, posso dizer com novos exemplos de músicas se esta pessoa gostaria ou não da música em questão que está sendo testada.

Nesse passo, Marko Kolanovic e Rajesh T. Krishnamachari⁷¹ colaboram ao dizer que aprendizado de máquina supervisionado é o caso em que um algoritmo de

⁷⁰ PEIXOTO, Fabiano Hartmann; SILVA, Roberta Zumblick Martins. **Inteligência Artificial e Direito**. Curitiba: Alteridade Editora, 2019, p. 92-93.

⁷¹ KOLANOVIC, Marko; KRISHNAMACHARI, Rajesh T. **Big Data and AI Strategies: Machine Learning and Alternative Data Approach to Investing**. 2017. Disponível em: <https://cpb-us-e2.wpmucdn.com/faculty.sites.uci.edu/dist/2/51/files/2018/05/JPM-2017-MachineLearningInvestments.pdf>. Acesso em: 28 ago. 2021.

machine learning calibra seus parâmetros usando dados alimentados por um analista, a partir disso o algoritmo ajusta o modelo e o melhora à medida que mais dados são coletados. O termo “aprendizado supervisionado” surge a partir da observação de que o analista guia, ou seja, supervisiona o algoritmo na calibração de parâmetros, fornecendo um conjunto de variáveis de entrada e saída claramente rotuladas ou variáveis previstas. Um exemplo de utilização desse tipo de aprendizado seria tentar prever o quanto o mercado irá se mover se houver um aumento repentino na inflação.

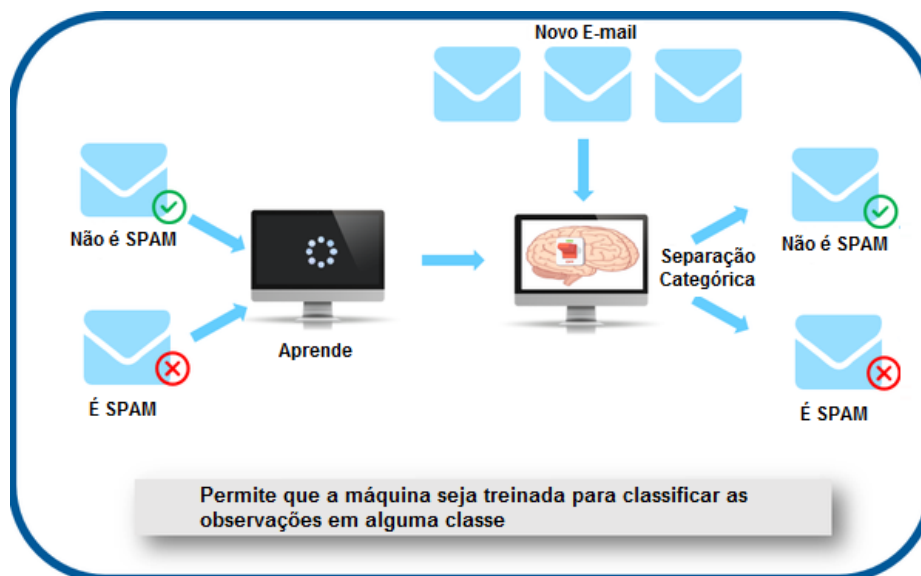
Portanto, o algoritmo supervisionado busca através do espaço de hipóteses possíveis (classes) por aquele que terá melhor desempenho para o dado, isso acontece também para conjuntos que não fazem parte do conjunto de treinamento. Com um conjunto de testes de exemplos que são distintos do conjunto de treinamento ou com novos exemplos com diferentes valores, podemos medir a precisão de uma hipótese de um modelo gerado. Dizemos que uma hipótese generaliza bem se prevê corretamente o valor de saída para novos exemplos.⁷²

Por fim, Goodfellow, Bengio e Courville⁷³ salientam que os aprendizados supervisionados e não supervisionados não possuem uma definição formal, e que as linhas que os separam não são claras. Muitas tecnologias de *machine learning* podem ser usadas para desempenhar nas duas formas. Ainda que não se trate de conceitos formais completamente distintos, eles ajudam a categorizar de forma aproximada um pouco do que se faz com algoritmos de aprendizado de máquina.

Figura 2 – Exemplo de aprendizagem supervisionada

⁷² RUSSELL, Stuart; NORVIG, Peter. **Inteligência Artificial**. 3. ed., Rio de Janeiro: Editora Campus-Elsevier, 2013.

⁷³ GOODFELLOW, Ian; BENGIO, Yoshua; COURVILLE, Aaron. **Deep Learning**. Disponível em: https://www.deeplearningbook.org/front_matter.pdf. Acesso em: 18 ago. 2021.



Fonte: Walter V. Clavera⁷⁴

Como mostrado na figura acima, inicialmente se pega alguns dados e os marca como 'Spam' ou 'Não é Spam'. Esses dados rotulados são usados pelo modelo de treinamento supervisionado, esses dados são usados para treinar o modelo. Uma vez treinado, pode-se testar o modelo experimentando-o com alguns novos e-mails de teste e a verificação do modelo é capaz de prever a saída correta.

1.3.1.2 Aprendizagem não supervisionada

No que tange ao aprendizado não supervisionado esse se assemelha mais com a forma que os seres humanos aprendem. Não há uma introdução prévia sobre as classificações e padrões de inteligência. É um problema menos definido, uma vez que não se sabe o tipo de padrão que se espera. Neste sentido, há uma maior abrangência de aplicação do que o aprendizado supervisionado, pois não necessita de uma introdução prévia e humana.⁷⁵

Nesse sentido, aprendizado de máquina não supervisionado é definido como um sistema que aprende sem precisar de assistência humana e dados "rotulados" para ensiná-lo, ou seja, o algoritmo procura encontrar estrutura nos dados, geralmente agrupando observações. Em finanças esse método de aprendizado pode ser usado

⁷⁴ CLAVERA, Walter V. Aprendizado de Máquina (Machine Learning). *In: Sejatech e Falettech*, 2019. Disponível em: <https://www.redesdaude.com.br/aprendizado-de-maquina-machine-learning/>. Acesso em: 28 ago. 2021.

⁷⁵ PEIXOTO, Fabiano Hartmann; SILVA, Roberta Zumblick Martins. *Inteligência Artificial e Direito*. Curitiba: Alteridade Editora, 2019.

na identificação de regimes históricos como regime de alta ou baixa volatilidade, regime de inflação crescente, entre outros, importante para alocação entre diferentes ativos e prêmios de risco.⁷⁶

No aprendizado não supervisionado não se utiliza valores de referência, isto é, não existe um treinamento pré-programado para que o ambiente adquira conhecimento. Lorena e Carvalho⁷⁷ ressaltam que no ambiente não supervisionado, o algoritmo “aprende a representar (ou agrupar) as entradas submetidas, segundo medidas de similaridade”. As técnicas de aprendizado não supervisionado são mais usadas quando a compreensão dos dados é feita através de padrões ou tendências.

Nos dados de entrada não se apresenta um ponto de saída, ponto de conclusão, não tem uma “resposta certa” rotulada de saída. Cabe ao algoritmo descobrir semelhanças entre os dados e agrupá-los adequadamente. O objetivo é explorar os dados e encontrar alguma estrutura dentro deles. O aprendizado não-supervisionado funciona bem com dados transacionais. Por exemplo, ele pode identificar segmentos de clientes com atributos similares que podem, então, ser tratados de modo igualmente em campanhas de marketing. Algoritmos populares utilizados incluem mapas auto organizáveis, mapeamento por proximidade, agrupamento K-means. Esses algoritmos também são utilizados para segmentar tópicos de texto, recomendar itens e identificar pontos discrepantes nos dados.

Kevin Murphy⁷⁸ comenta que o aprendizado não supervisionado também costuma ser chamado de descoberta do conhecimento. É um problema menos definido, pois não se sabe que tipos de padrões se está procurando e não há uma métrica clara de erro para controle – ao contrário do aprendizado supervisionado, em que se sabe o tipo de padrão almejado e a métrica adotada.

O aprendizado não supervisionado tem o objetivo de proceder segmentações de grupo e dispensam a rotulação, sendo muito utilizado em sistemas de recomendação, separação de clientela por gostos e afinidades, entre outros casos. Os algoritmos não supervisionados não apresentam um padrão já estipulado de

⁷⁶ KOLANOVIC, Marko; KRISHNAMACHARI, Rajesh T. **Big Data and AI Strategies: Machine Learning and Alternative Data Approach to Investing**. 2017. Disponível em: <https://cpb-us-e2.wpmucdn.com/faculty.sites.uci.edu/dist/2/51/files/2018/05/JPM-2017-MachineLearningInvestments.pdf>. Acesso em: 28 ago. 2021.

⁷⁷ LORENA, Ana Carolina; CARVALHO, André C. P. L. F. de. Uma introdução às Support Vector Machines. *In: Revista de Informática Teórica e Aplicada*. vol. 14, n.2, 2007. Disponível em: https://seer.ufrgs.br/rita/article/download/rita_v14_n2_p43-67/3543. Acesso em: 28 ago. 2021, p. 44.

⁷⁸ MURPHY, Kevin P. **Machine Learning: A Probabilistic Perspective**. 2012. Disponível em: <https://www.cs.ubc.ca/~murphyk/MLbook/pml-intro-22may12.pdf>. Acesso em: 28 ago. 2021.

informações e resultados, cabendo ao próprio sistema fazer todo o trabalho, havendo ausência do esforço humano em sua execução. Logo, há uma autonomia maior da máquina.

Figura 3 – Exemplo de aprendizagem não supervisionada



Fonte: Walter V. Clavera⁷⁹

Na figura acima, se dá algumas características ao modelo do que são ‘Patos’ e ‘Não são Patos’. Nos dados de treinamento, não se fornece nenhum rótulo aos dados correspondentes. O modelo não supervisionado é capaz de separar os dois caracteres observando o tipo de dados e modelando a estrutura ou distribuição subjacente nos dados para aprender mais sobre eles. Após analisar os dados, aprende a representar as entradas submetidas, e assim, reconhecer quando são patos ou não segundo medidas de similaridade.

1.3.1.3 Aprendizagem por reforço

Finalmente, de acordo com os autores Roberta Zumblick Martins da Silva e Fabiano Hartmann Peixoto⁸⁰, a aprendizagem por reforço é a menos utilizada e conhecida, ela é útil para se aprender como agir ou se comportar por sinais de recompensa e punição. Basicamente trata de uma máquina que pensa como os pesquisadores, a partir dela é avaliada a efetividade de métodos, analisando meios de resolver problemas de interesse científico e econômico através de diagnósticos

⁷⁹ CLAVERA, Walter V. Aprendizado de Máquina (Machine Learning). *In: Sejatech e Falettech*, 2019. Disponível em: <https://www.redesdaude.com.br/aprendizado-de-maquina-machine-learning/>. Acesso em: 28 ago. 2021.

⁸⁰ PEIXOTO, Fabiano Hartmann; SILVA, Roberta Zumblick Martins. *Inteligência Artificial e Direito*. Curitiba: Alteridade Editora, 2019.

matemáticos e experimentos computacionais.

Aprendizado por reforço é o caso em que a conexão entre ações e recompensas é desconhecida e deve ser inferida a partir das interações do agente com o ambiente. Um conjunto de dados não rotulados é enviado a um algoritmo e, em seguida, o algoritmo escolhe uma ação para cada ponto de dados. Um humano pode então fornecer feedback (a parte supervisionada) que subsequentemente ajuda o algoritmo a aprender.⁸¹

Aprendizagem por reforço é um método no qual o agente, sem o conhecimento prévio de quais ações são melhores de serem executadas, deve experimentar o ambiente e aprender conforme interage com ele. A expectativa é que, conforme o agente interage com o ambiente, seja capaz de melhorar seu desempenho. É uma área da aprendizagem de máquina inspirada pela psicologia comportamental, sendo comumente chamada de método da tentativa e erro, simulando a forma de aprendizagem humana. Em problemas de aprendizagem por reforço, o ambiente e as regras que o regem são desconhecidos pelo agente. Nesta forma de aprendizagem, o agente deve aprender a partir de suas próprias experiências. Não há uma classificação específica de ações corretas ou erradas, apenas recompensas que indicam se o agente está mais próximo ou longe do seu objetivo.⁸²

É comum haver confusão ao diferenciar o aprendizado por reforço com o aprendizado não supervisionado. Diante disto, Sutton e Barto⁸³ sustentam que a aprendizagem por reforço difere do aprendizado não supervisionado, que tipicamente se resume em encontrar estruturas escondidas em conjunto de dados não rotulados. Afirmam que os termos supervisionado e não supervisionado podem parecer exaustivos para a classificação dos paradigmas de *machine learning*, mas não são. Ainda que a ideia de se pensar que a aprendizagem por reforço seja um tipo de aprendizagem não supervisionado por não contar com exemplos do comportamento esperado, o aprendizado por reforço busca maximizar um sinal de bom desempenho, não encontrar uma estrutura oculta. Essa função, encontrar uma estrutura oculta,

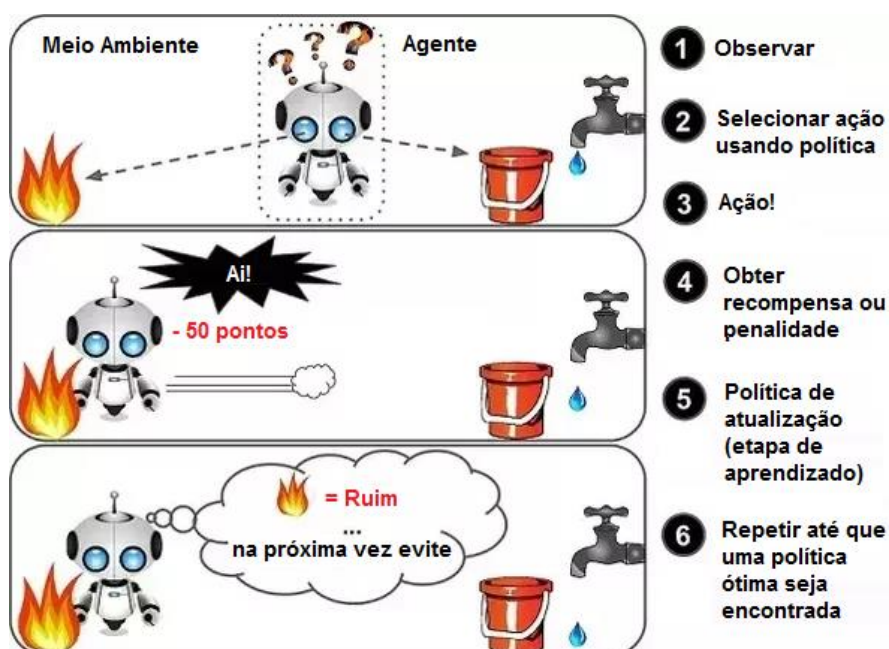
⁸¹ BUCHANAN, Bonnie G. **Artificial intelligence in Finance**. Washington: Zenodo, 2019. Disponível em: <https://zenodo.org/record/2612537#.YUaPM7hKjIV>. Acesso em: 28 ago. 2021.

⁸² SUTTON, Richard S.; BARTO, Andrew G. **Reinforcement learning: An introduction**. 2. ed., London: MIT press, 1998.

⁸³ SUTTON, Richard S.; BARTO, Andrew G. **Reinforcement learning: An introduction**. 2. ed., London: MIT press, 1998.

pode até se mostrar útil no aprendizado por reforço, mas por si só não dá conta do problema de maximizar um sinal de recompensa. Consideram, portanto, a aprendizagem por reforço um terceiro paradigma da aprendizagem de máquina.

Figura 4 – Exemplo de aprendizagem por reforço



Fonte: Walter V. Clavera⁸⁴

Na figura acima, pode-se observar que o agente recebe 2 opções, isto é, um caminho com água ou um caminho com fogo. Um algoritmo de reforço funciona na recompensa de um sistema, ou seja, se o agente usa o caminho do fogo, então as recompensas são subtraídas e o agente tenta aprender que deve evitar o caminho do fogo. Se tivesse escolhido o caminho da água ou o caminho seguro, alguns pontos teriam sido adicionados aos pontos de recompensa, o agente então tentaria descobrir que caminho é seguro e qual caminho não é. É basicamente aproveitando as recompensas obtidas, que o agente melhora o conhecimento do ambiente para selecionar a próxima ação.

1.3.2 Deep Learning

Já o termo *deep learning*, remete a uma subcategoria de *machine learning*,

⁸⁴ CLAVERA, Walter V. Aprendizado de Máquina (Machine Learning). *In: Sejatech e Falettech*, 2019. Disponível em: <https://www.redesdaude.com.br/aprendizado-de-maquina-machine-learning/>. Acesso em: 28 ago. 2021.

que diz respeito a oportunidades de aprendizagens profundas com o uso de redes neurais para incrementar resultados da inteligência em diversas esferas, como reconhecimento facial, reconhecimento vocal, dentre outros.⁸⁵

Define-se, segundo Russell e Norvig⁸⁶, como um tipo de sistema computacional inspirado pelas propriedades básicas de neurônios biológicos. Sendo elas compostas por muitas unidades individuais conectadas por ligações direcionadas.

Dessa forma, este mecanismo é composto por camadas. Ao cortamos uma imagem qualquer em milhares de pedaços e enviarmos para a primeira camada do deep learning, essa receberá a imagem e analisará, realizando seu trabalho, em seguida essa imagem analisada pela primeira camada, passará para a segunda, fazendo sua respectiva leitura, e assim por diante até que a camada final consiga produzir a solução do problema.

Ato contínuo, cada neurônio deposita seu peso para os dados que entram, os chamados *inputs*, e mandam para os que saem em direção a próxima camada, os chamados *outputs*. O resultado então é determinado pelo peso depositado por cada neurônio. Desta forma, o sistema pode adquirir maior potência na medida que se acrescenta camadas e unidades a rede.

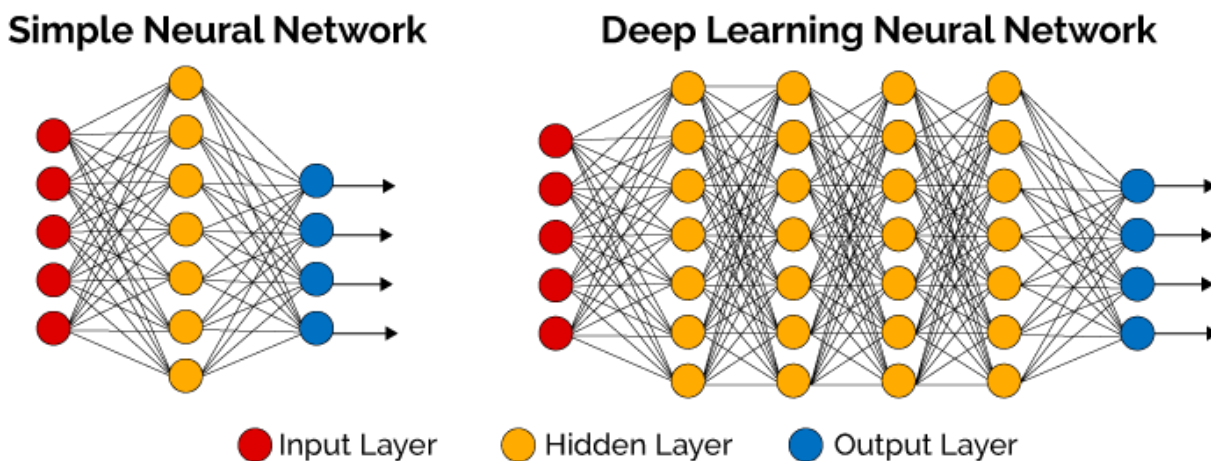
O que diferencia *Deep Learning* de *Machine Learning* é que o Aprendizado de Máquina (*Machine Learning*) funciona de uma forma linear, enquanto o *Deep Learning* se organiza e funciona em camadas encadeadas, o que possibilita análises mais complexas e profundas com uma quantidade maior de dados. Usualmente essas camadas são classificadas em três grupos: as de entrada (input layer, em inglês), onde os dados são apresentados, que fazem conexões com as camadas intermediárias ou escondidas (hidden layers, em inglês), na qual acontece a maior parte do processamento, e, por fim, a camada de saída (output layer, em inglês), onde o resultado final é concluído e apresentado.⁸⁷ A imagem abaixo ilustra a atividade mencionada:

⁸⁵ MAGRANI, Eduardo. **Entre dados e robôs: ética e privacidade na era da hiperconectividade**. 2. ed., Porto Alegre: Arquipélogo Editorial, 2019.

⁸⁶ RUSSELL, Stuart; NORVIG, Peter. **Artificial intelligence: a modern approach**. New Jersey: Editora Pearson Prentice Hall, 2010.

⁸⁷ CARVALHO, André. **Redes Neurais Artificiais**. Disponível em: <https://sites.icmc.usp.br/andre/research/neural/>. Acesso em: 28 ago. 2021.

Figura 5 – Demonstração Deep Learning



Fonte: Laboratório iMobilis⁸⁸

Sendo assim, as camadas em uma rede neural profunda correspondem a níveis de abstração ou composição. O número de camadas e o tamanho de cada uma impacta diretamente no nível de abstrações que a rede é capaz de fornecer. As camadas de nível superior aprendem com as camadas de nível inferior, trabalhando em uma abstração acima da camada prévia. Desta forma, uma rede treinada com um espaço amostral grande e diverso o suficiente, consegue fornecer uma saída correta à uma entrada até então nunca vista, por ter aprendido a reconhecer padrões graças aos dados de treinamento. Técnicas de *Deep Learning* permitem que se lide com uma quantidade alta de dados e com dados de abstrações mais complexas. Por isso, é comum que várias arquiteturas de *Deep Learning* sejam utilizadas nos campos da visão computacional, do reconhecimento de fala e do processamento de linguagem natural.⁸⁹

São exemplos de *deep learning*: carros que ligam sozinhos, reconhecimento faciais e de objetos em fotos, vídeos e compreensão de linguagem em traduções.

A *deep learning*, portanto, vai além da perspectiva neurocientífica sobre a atual geração de modelos de aprendizagem de máquinas. Apela a um princípio mais geral de aprendizagem de múltiplos níveis de composição, que podem ser aplicados

⁸⁸ COELHO, Mateus. Inteligência Artificial e Deep Learning. *In: Laboratório iMobilis*, 2019. Disponível em: <http://www2.decom.ufop.br/imobilis/inteligencia-artificial-e-deep-learning/>. Acesso em: 30 jul. 2021.

⁸⁹ BENGIO, Yoshua; COURVILLE, Aaron; VINCENT, Pascal. Representation learning: A review and new perspectives. *In: IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 35, n. 8, p. 1798–1828, 2013.

em estruturas de aprendizagem de máquinas que não são necessariamente inspiradas nos sistemas neurais humanos, consoante explica Eduardo Magrani⁹⁰.

1.3.3 Redes Neurais Artificiais

As Redes Neurais Artificiais (RNAs) são modelos de aprendizagem de máquina baseados na atividade das redes de neurônios biológicos, assim como as do cérebro humano. Existe um estímulo em pesquisar a forma de como o cérebro processa as informações por este ser complexo, não-linear e paralelo. No entanto, cabe salientar que o cérebro humano tem a capacidade de organizar seus constituintes estruturais (neurônios) de forma a realizar certos processamentos (reconhecimento de padrões, percepção e controle motor, entre outros) de maneira muito mais rápida e eficiente que um supercomputador.⁹¹

As Redes Neurais Artificiais são provavelmente a mais antiga técnica de IA. Surgiu na década de 40, por Walter Pitts e McCulloch, o primeiro matemático e o segundo neurofisiologista. A ideia era fazer uma analogia entre neurônios biológicos e circuitos eletrônicos, capazes de simular conexões sinápticas pelo uso de resistores variáveis e amplificadores. Foi ainda nesta década de 40, que o matemático Johann Von Neumann, propôs a arquitetura dos computadores eletrônicos atuais, em seu trabalho intitulado “Preliminary Discussion of the Logic Design of an Electronic Computing Instrument”. Uma máquina de processamento sequencial com CPU e memória separados e um ponteiro que registra o endereço do próximo comando a ser executado.⁹²

Segundo Haykin⁹³, as RNAs se assemelham ao cérebro em dois aspectos: 1) a partir do ambiente, a informação é transmitida pela rede que, por sua vez, passa por um processo de aprendizagem, resultando em um conhecimento adquirido; 2) as forças de conexão entre neurônios, também chamadas de pesos sinápticos, são utilizadas para armazenar o conhecimento adquirido. Os neurônios artificiais, por sua

⁹⁰ MAGRANI, Eduardo. **Entre dados e robôs: ética e privacidade na era da hiperconectividade**. 2. ed., Porto Alegre: Arquipélogo Editorial, 2019.

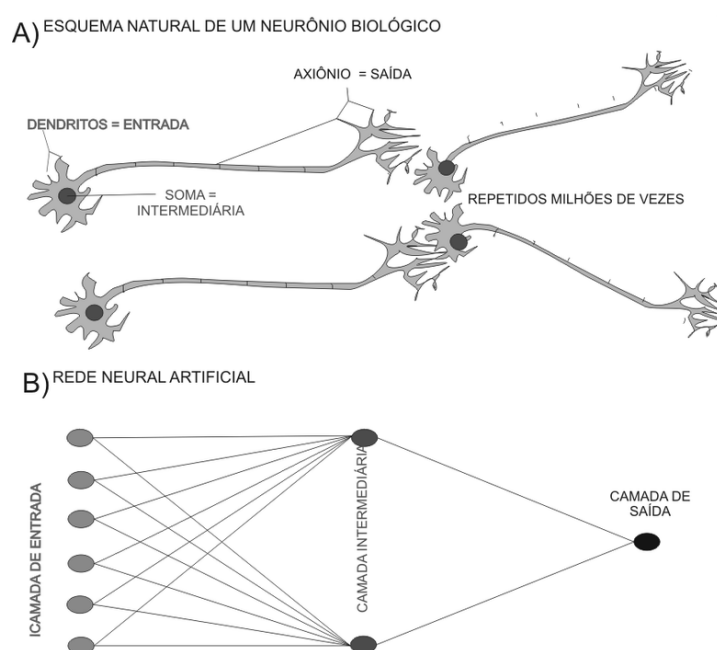
⁹¹ HAYKIN, Simon. **Redes Neurais: princípios e prática**. 2. ed., Porto Alegre: Bookman Companhia, 2001.

⁹² LUDWIG, Oswaldo Jr.; MONTGOMERY, Eduard. **Redes Neurais: Fundamentos e Aplicações em C**. 1. ed., Rio de Janeiro: Editora Ciência Moderna, 2007.

⁹³ HAYKIN, Simon. **Redes Neurais: princípios e prática**. 2. ed., Porto Alegre: Bookman Companhia, 2001.

vez, são compostos por três elementos básicos: (1) as sinapses, responsáveis por propagar as informações entre neurônios, cada uma possuindo um peso próprio; (2) um somador, que soma as informações de entrada providas das sinapses e geram uma nova informação; (3) uma função de ativação, responsável por restringir a amplitude do sinal produzido pelo neurônio em um valor finito permissível para propagar a informação (sinapse).

Figura 6 - Estrutura de um neurônio biológico comparado com uma rede neural artificial



Fonte: Rafael Manica⁹⁴

Os neurônios artificiais, também chamados de nós, são as unidades que compõem uma rede neural artificial. Estes neurônios recebem uma série de valores de entrada, juntamente com um peso, que determina a intensidade e o sinal da conexão. Cada neurônio possui uma função de ativação, que recebe a soma de todas as entradas multiplicadas pelos seus respectivos pesos, processa tais valores e gera uma saída para o neurônio.⁹⁵

⁹⁴ MANICA, Rafael. Aplicação de uma rede neural artificial simplificada para a identificação de gradação de depósitos turbidítico. *In: Geociências*, vol. 32, n. 3, p. 429-440, 2013. Disponível em: https://www.researchgate.net/publication/260201326_APLICACAO_DE_UMA_REDE_NEURAL_ARTIFICIAL_SIMPLIFICADA_PARA_A_IDENTIFICACAO_DE_GRADACAO_DE_DEPOSITOS_TURBIDITICO. Acesso em: 29 ago. 2021.

⁹⁵ RUSSELL, Stuart; NORVIG, Peter. *Inteligência Artificial*. 3. ed., Rio de Janeiro: Editora Campus-Elsevier, 2013.

Como se pode observar, a rede neural artificial é parte essencial para o *machine learning*. A arquitetura de uma RNA é composta por camadas de neurônios, possuindo, no mínimo, duas camadas: (1) a camada de entrada (input) que analisa e processa as informações dos conjuntos de dados; (2) camada de classificação (output), que através dos dados processados na(s) camada(s) prévia(s), separa-os em classes com base nos padrões característicos previamente conhecidos. Esses outputs são então comparados com os valores esperados obtidos através dos exemplos de treinamento e a taxa de erro é calculada para então atualizar os pesos das sinapses, visando diminuir a taxa de erro. RNAs podem contar com outras camadas entre a de entrada e a de classificação, denominadas camadas intermediárias. Quando essas RNAs possuem várias camadas (normalmente mais de 20), usa-se a expressão *deep learning* para se referir a essas redes. Diversos tipos de RNAs do tipo *deep learning* têm sido desenvolvidas nos últimos anos para os mais diversos fins, incluindo a detecção automática de objetos, segmentação em imagens e classificação de doenças.⁹⁶

Vale ressaltar, ainda, que mesmo que ainda não se tenha conseguido construir uma rede neural artificial com as mesmas capacidades do cérebro humano, de forma geral os estudos das redes neurais ampliam o conhecimento teórico, aprimorando a compreensão do funcionamento do sistema nervoso, assim como possibilitam a melhoria de aplicações fundamentais no cotidiano da população, a exemplo da neurofarmacologia, que simula o efeito de drogas sobre o sistema nervoso central. Outra vantagem da neurociência computacional é a redução do uso de animais de laboratório e a validação de experimentos biológicos. Por fim, pesquisas interdisciplinares que conectam a neurociência com demais práticas sistemáticas, tais como a computação, são de grande valia e estão em constante aperfeiçoamento no campo da neurociência.

⁹⁶ HAYKIN, Simon. **Redes Neurais: princípios e prática**. 2. ed., Porto Alegre: Bookman Companhia, 2001.

2 A APLICAÇÃO DA INTELIGÊNCIA ARTIFICIAL NO DIREITO E A SUA ATUAL REGULAMENTAÇÃO

A onipresença da tecnologia é incontestável, atualmente ela vem ganhando espaço em cada um dos cenários da vida cotidiana, nas residências, nas indústrias, nos escritórios e, inclusive nos tribunais, com isto, nota-se a crescente intensidade desse sistema acompanhando nossas vidas por meio de algoritmos conectados à internet. Essa introdução tecnológica alterou de forma profunda nosso relacionamento com o mundo, principalmente, em relação as outras pessoas, ao trabalho e ao acesso às informações.

No ambiente jurídico a tecnologia da inteligência artificial vem ganhando seu espaço, e apesar de ainda haverem muitos questionamentos sobre qual o verdadeiro potencial que essas máquinas irão assumir dentro do judiciário, é importante mencionar que atualmente, as máquinas estão atuando apenas com a capacidade de aprendizado supervisionado. Diante disso, é possível vislumbrar algumas das principais aplicações da inteligência artificial no Direito: a previsão de resultados litigiosos; análise de documentos; pesquisa de jurisprudência e revisão contratual; identificação de padrões de decisões judiciais; identificação de propriedade intelectual em portfólios; e faturamento automático de honorários, com isto, resta evidente que a inteligência artificial ainda não está voltada para tomar decisões sozinha, ela apenas reconhece padrões previamente estruturados pelo ser humano que a supervisiona.⁹⁷ Para que se possa vislumbrar melhor a implementação desta tecnologia no judiciário, este capítulo se propõe a estudar as principais ferramentas da IA utilizadas no âmbito jurídico. Além disto, se apresentará no presente capítulo a atual regulamentação no ordenamento jurídico sobre a IA.

2.1 PRINCIPAIS FERRAMENTAS DA INTELIGÊNCIA ARTIFICIAL UTILIZADAS NO ÂMBITO JURÍDICO EM PÂRAMETRO INTERNACIONAL

A adoção e aplicação de soluções através de computadores, processamentos digitais e máquinas que aprendem e executam tarefas relativamente complicadas de

⁹⁷ PINTO, Henrique Alves. A utilização da inteligência artificial no processo de tomada de decisões. **In: RIL Brasília**, ano 57, n. 225, p. 43-60, jan./mar. 2020. Disponível em: https://www12.senado.leg.br/ril/edicoes/57/225/ril_v57_n225_p43.pdf. Acesso em: 30 jul. 2021.

forma precisa e com uma rapidez indiscutível ao desempenho humano toma uma evidente notoriedade no mundo inteiro. Sendo assim, é importante analisar algumas das principais ferramentas utilizadas no judiciário, para que se possa vislumbrar como a tecnologia da inteligência artificial vem sendo utilizada, além do mais, ressalta-se que as ferramentas apresentadas não são exaustivas, são apenas exemplificativas, tendo em vista que o sistema jurídico tem utilizado da Inteligência Artificial para auxiliar em diversas atividades, sendo que se optou por estas por serem as mais conhecidas até o momento.

2.1.1 Sistema COMPAS

Os Estados Unidos já vêm utilizando um sistema denominado COMPAS (Correctional Offender Management Profiling for Alternative Sanctions), este é um software privado que utiliza algoritmos matemáticos para determinar o risco de reincidência de alguns acusados na área criminal, auxiliando a decisão dos magistrados em relação à prisão preventiva. Os níveis variam entre alto, médio e baixo grau de reincidência.⁹⁸

Todavia, este sistema tem sido criticado em decorrência de uma pesquisa realizada pela ONG ProPublica, reportada pela BBC americana em 2016 e que se tornou história documentada pela plataforma Netflix, denominada “Making a Murderer”. Neste estudo foi concluído que este sistema possui julgamentos enviesados, apresentando em suas decisões vieses preconceituosos perante diferentes etnias apresentadas ao sistema.⁹⁹ Nesta perspectiva, Julia Angwin, pesquisadora da ONG, afirma em entrevista a Simon Maybin jornalista da BBC News¹⁰⁰:

Quando analisamos um acusado negro e outro branco com a mesma idade, sexo e ficha criminal - e levando em conta que depois de serem avaliados os dois cometeram quatro, dois ou nenhum crime -, o negro tem 45% mais

⁹⁸ LAGE, Lorena Muniz e Castro; LIMA, Henrique Cunha Souza; MARTINO, Antonio Anselmo. Inteligência Artificial e Tecnologias Aplicadas ao Direito – II. **In: Congresso Internacional de Direito e Inteligência Artificial**. Belo Horizonte: Skema Business School, 2020. Disponível em: <https://www.conpedi.org.br/wp-content/uploads/2020/09/SKEMA-Intelig%C3%A2ncia-Artificial-e-tecnologias-aplicadas-ao-Direito-II.pdf>. Acesso em: 02 ago. 2021.

⁹⁹ MAYBIN, Simon. Sistema de algoritmo que determina pena de condenados criapolêmica nos EUA. **In: BBC News**. 2016. Disponível em: <https://www.bbc.com/portuguese/brasil-37677421>. Acesso em: 02 ago. 2021.

¹⁰⁰ MAYBIN, Simon. Sistema de algoritmo que determina pena de condenados criapolêmica nos EUA. **In: BBC News**. 2016. Disponível em: <https://www.bbc.com/portuguese/brasil-37677421>. Acesso em: 02 ago. 2021.

chances do que o branco de receber uma pontuação alta.

A sistemática consiste em um questionário inserido em um sistema de pontos de um a dez, buscam identificar questões criminais na família, nos ciclos de amizades e histórico escolar e profissional. Ao final, são feitas perguntas estratégicas para identificar pensamentos voltados ao crime. A partir deste resultado, o magistrado decide, com base no score obtido, se o suspeito poderá ser solto através de fiança, deverá seguir para um presídio, ou, ainda, se deverá receber um tipo específico de sentença, como a concessão da liberdade condicional.¹⁰¹

Esta situação envolvendo vieses discriminatórios ficou evidente no caso, abordado pelos pesquisadores, envolvendo Vernon Prater, sendo pego roubando uma loja que gerou um prejuízo próximo a U\$ 86,00. O indivíduo já havia sido condenado por um assalto à mão armada e outras diversas infrações penais. E, do outro lado, Brisha Borden, que atrasada para determinado compromisso, pegou uma scooter que pertencia a uma criança para que chegasse ao local a tempo. Logo em seguida, foi flagrada e conseqüentemente levada a prisão por assalto.¹⁰²

Dito isso, ambos os acusados foram submetidos ao questionário do sistema COMPAS. Brisha, mulher negra, foi classificada como de alto risco, obteve score de número 8. Já Vernon, homem branco, foi classificado como de baixo risco com score de número 3.

Outra questão relevante quanto a discriminação do sistema é a ficha criminal de Vernon, que consiste em 2 assaltos a mão armada; 1 tentativa de assalto a mão armada e um grande roubo. A de Brisha, constam 4 delitos juvenis, que constituem delitos banais.¹⁰³

Analisando outros casos, a reportagem demonstrou que o programa sistematicamente atribuí a negros probabilidades maiores de cometer crimes, ainda que eles tivessem histórico e perfil similares ao de brancos.¹⁰⁴

¹⁰¹ STEIN, Raphael Wilson L. É Preciso Evolução Na Inteligência Artificial? *In: JusBrasil*. 2019. Disponível em: <https://stein.jusbrasil.com.br/artigos/743864200/e-preciso-evolucao-na-inteligencia-artificial>. Acesso em: 02 ago. 2021.

¹⁰² STEIN, Raphael Wilson L. É Preciso Evolução Na Inteligência Artificial? *In: JusBrasil*. 2019. Disponível em: <https://stein.jusbrasil.com.br/artigos/743864200/e-preciso-evolucao-na-inteligencia-artificial>. Acesso em: 02 ago. 2021.

¹⁰³ STEIN, Raphael Wilson L. É Preciso Evolução Na Inteligência Artificial? *In: JusBrasil*. 2019. Disponível em: <https://stein.jusbrasil.com.br/artigos/743864200/e-preciso-evolucao-na-inteligencia-artificial>. Acesso em: 02 ago. 2021.

¹⁰⁴ STEIN, Raphael Wilson L. É Preciso Evolução Na Inteligência Artificial? *In: JusBrasil*. 2019. Disponível em: <https://stein.jusbrasil.com.br/artigos/743864200/e-preciso-evolucao-na-inteligencia-artificial>. Acesso em: 02 ago. 2021.

A ProPublica afirmou que o software da empresa prevê acertadamente 61% dos casos, porém os negros são duas vezes mais inclinados a reincidir e serem qualificados como risco maior. Inclusive, há casos em que juízes citaram as avaliações de risco e pontuações nas decisões, por exemplo, Zilly roubou objetos para manter seu vício em metanfetamina, porém estava se esforçando para se recuperar. Zilly foi apontado como tendo alto risco de reincidência e foi mandado para a prisão. Houve recurso e Brennan (criador do software) foi depor declarando que o software não tinha esse intuito de ser utilizado em sentenças, o objetivo era somente reduzir os crimes. A questão é que o COMPAS não deveria ser utilizado como meio de prova. A partir dessa declaração o juiz reduziu drasticamente a pena de Zill.¹⁰⁵

No entanto, a maior crítica ao software é de que o cálculo realizado pelo algoritmo não é público. O réu, e até mesmo o próprio juiz, não têm acesso ao como o resultado é alcançado, restringindo-se, dessa forma, a mostrar o número final da verificação que faz dentro da escala que utiliza. Isso se deve ao fato de que o sistema é de uma empresa privada, a qual se recusa a revelar o método utilizado em função deste se tratar de um segredo comercial.¹⁰⁶

A partir disto, a Suprema Corte dos Estados Unidos definiu que as avaliações de riscos devem ser acompanhadas de uma explicação contendo suas limitações e indicando o meio para se obter suas escolhas.¹⁰⁷

2.1.2 Robô ROSS

Com a evidente morosidade do sistema judicial devido ao alto fluxo de transmissão de informações, os escritórios de advocacia viram uma necessidade de ajuda quanto a suas tarefas. Em novembro de 2017, um dos mais conceituados escritórios de advocacia de falência de Nova York contratou o primeiro advogado robô

¹⁰⁵ ANGWIN, Julia; LARSON, Jeff; MATTU, Surya; KIRCHNER, Lauren. Machine Bias: There's software used across the country to predict future criminals. And it's biased against blacks. *In: ProPublica*, 2016. Disponível em: <https://www.propublica.org/article/machine-bias-risk-assessments-in-criminal-sentencing>. Acesso em: 02 ago. 2021.

¹⁰⁶ COUTO, João Victor Lima de Abreu. Desafios da inteligência artificial na interpretação jurídica em loomis vs. Wisconsin: novas tecnologias na interpretação de decisões judiciais sob uma ótica da hermenêutica de hans george gadamer. *In: Congresso Internacional de Direito e Inteligência Artificial - Inteligência artificial e tecnologias aplicadas ao direito II*. Belo Horizonte: Skema Business School, 2020, p. 43.

¹⁰⁷ STEIN, Raphael Wilson L. É Preciso Evolução Na Inteligência Artificial? *In: JusBrasil*. 2019. Disponível em: <https://stein.jusbrasil.com.br/artigos/743864200/e-preciso-evolucao-na-inteligencia-artificial>. Acesso em: 02 ago. 2021.

do mundo. Denominado como ROSS, o advogado robô é um colaborador cognitivo, com capacidade de aprender, pensar e tomar decisões de forma autônoma, pertencente a IA Forte, desenvolvido pela IBM, na cidade de Toronto no Canadá, 2016, com a Tecnologia Watson.¹⁰⁸

O ROSS foi desenvolvido com a intenção de auxiliar na área jurídica, sendo capaz de analisar um bilhão de documentos em um segundo e apresentar ao usuário exatamente aquilo que lhe interessa. Porém, é importante salientar que sua finalidade não está relacionada apenas em ser um mecanismo de pesquisa, ele possui a capacidade de ler, entender e interpretar problemas jurídicos, bem como apontar possíveis soluções viáveis.¹⁰⁹

Portanto, a plataforma é capaz de extrair conclusões ao analisar a literatura jurídica, selecionar informações relevantes para um caso específico, formular hipóteses, gerar respostas sustentadas por referências e interagir com o usuário. A interface do sistema é simples e intuitiva: o advogado faz uma pergunta e o robô soluciona a questão, citando precedentes jurídicos, leis relacionadas e até um percentual de confiabilidade da resposta fornecida. O sistema também é capaz de pesquisar em outros idiomas e alertar o advogado para novas mudanças de entendimento e tendências jurisprudenciais.¹¹⁰

Conforme explica Marcos Coelho¹¹¹, o sistema ROSS funciona em quatro etapas:

I) Análise de proposição: de forma alheia aos métodos tradicionais de busca com sistemas de palavras-chave, ROSS é capaz de realizar um exame de elementos do período linguístico — sujeito, verbo, objeto — para alcançar diversas interpretações semânticas; II) Geração de Hipótese: para cada informação encontrada, o sistema busca em sua memória informações relacionadas, formulando hipóteses; III) Pontuação de Hipóteses e Evidências: trata-se da avaliação de cada hipótese formulada, através da coleta de evidências positivas e negativas em sua base de dados, atribuindo

¹⁰⁸ MELO, João Ozório de. **Escritório de Advocacia estreia primeiro “robô-advogado” nos EUA**. 2016. Disponível em: <https://www.conjur.com.br/2016-mai-16/escritorio-advocaciaestreiaprimero-robo-advogado-eua>. Acesso em: 02 ago. 2021.

¹⁰⁹ AMORIM, Fernanda Pacheco. **Respeita as Mina: Inteligência artificial e violência contra a mulher**. 1. ed., Florianópolis: EMais Editora, 2019, p. 102.

¹¹⁰ LANNES, Yuri Nathan da Costa; VALENTINI, Rômulo Soares; PIMENTA, Raquel Betty de Castro. **Inteligência Artificial e Tecnologias Aplicadas ao Direito – III. In: Congresso Internacional de Direito e Inteligência Artificial**. Belo Horizonte: Skema Business School, 2020. Disponível em: <https://www.conpedi.org.br/wp-content/uploads/2020/09/SKEMA-Intelig%C3%Aancia-Artificial-e-tecnologias-aplicadas-ao-Direito-III.pdf>. Acesso em: 02 ago. 2021, p. 29.

¹¹¹ COELHO, João Victor de Assis Brasil Ribeiro. **Aplicações e Implicações da Inteligência Artificial no Direito**. 2017. 61 f. TCC (Graduação em Direito) – Universidade de Brasília, Brasília, 2017. Disponível em: https://bdm.unb.br/bitstream/10483/18844/1/2017_JoaoVictordeAssisBrasilRibeiroCoelho.pdf. Acesso em: 02 ago. 2021, p. 32-33.

uma pontuação às possibilidades disponíveis e; IV) Comparação e Classificação: por último, após pontuar as hipóteses, ROSS utiliza sua experiência armazenada e algoritmos para definir os resultados com maior probabilidade de acerto.

Ainda de acordo com Marcos Coelho¹¹², o desempenho da plataforma é baseado em três pilares, com os seguintes parâmetros:

(i) Qualidade da Informação Provida - a completude da busca, a precisão nos resultados e sua relevância; (ii) Satisfação do Usuário – considerando a facilidade de utilização da plataforma e a confiança nela depositada; e (iii) Eficiência da Pesquisa – tempo necessário para que o usuário obtenha uma resposta satisfatória.

De acordo com a ROSS Intelligence¹¹³, esta tecnologia permite que os sistemas de IA tomem a linguagem cotidiana e as palavras faladas ou escritas e as transformem em conceitos e entidades, e dados estruturados. As vantagens em termos de potencial economia de tempo e custos são enormes.

Não obstante, a eficiência do robô vem sendo destacada por sua rapidez e precisão sendo requisitada em diversas áreas como a da medicina e da agricultura.

O robô ROSS não é o único que vem sendo aplicado no mundo jurídico, Jill Watson é um robô professor que auxilia estudantes de pós-graduação nos Estados Unidos a resolverem problemas em seus projetos tecnológicos.

2.1.3 LawGeex

O programa LawGeex é um software que utiliza de algoritmos de *Machine Learning* de Aprendizado Supervisionado e de análise de textos para ler peças contratuais, sendo capaz de fazer correções, alertas e sugerir alterações com base em políticas e princípios pré-estabelecidos por uma empresa ou escritório, automatizando processos de revisão e aprovação contratual. O programa foi iniciado em 2014 e já conta com mais de U\$ 9.5 milhões de investimento. Nas palavras de Noory Bechor, cofundador e CEO, “LawGeex está transformando operações legais cotidianas ao automatizar a revisão e aprovação de contratos. Nós ajudamos negócios a funcionarem de maneira mais rápida e suave ao remover o “engarrafamento

¹¹² COELHO, João Victor de Assis Brasil Ribeiro. **Aplicações e Implicações da Inteligência Artificial no Direito**. 2017. 61 f. TCC (Graduação em Direito) – Universidade de Brasília, Brasília, 2017. Disponível em: https://bdm.unb.br/bitstream/10483/18844/1/2017_JoaoVictordeAssisBrasilRibeiroCoelho.pdf. Acesso em: 02 ago. 2021, p. 33.

¹¹³ ROSS INTELLIGENCE. **Artificial Intelligence (AI) for the practice of law: An introduction**. 2018. Disponível em: <https://blog.rossintelligence.com/post/ai-introduction-law>. Acesso em: 02 ago. 2021.

jurídico”.¹¹⁴

Em um estudo divulgado pela própria LawGeex (Comparing the Performance of Artificial Intelligence to Human Lawyers in the Review of Standard Business Contracts), advogados americanos com décadas de experiência em direito societário e revisão de contratos foram confrontados com um computador para detectar problemas em cinco contratos NDA (NonDisclosure Agreement). Os profissionais humanos competiram contra um sistema LawGeex, que foi desenvolvido por três anos e treinado através de *machine learning* com base em dezenas de milhares de contratos. Após extensos testes, o sistema alcançou uma média de 94% de acertos na identificação de cláusulas problemáticas, enquanto os advogados atingiram um índice de 85%. Em média, foram necessários 92 minutos para que os profissionais humanos analisassem todos os cinco NDA’s propostos. O advogado que consumiu mais tempo gastou 156 minutos na análise, enquanto o profissional mais rápido fez a revisão em 51 minutos. Por sua vez, o computador concluiu a tarefa em apenas 26 segundos.¹¹⁵

2.2 PRINCIPAIS FERRAMENTAS DA INTELIGÊNCIA ARTIFICIAL UTILIZADAS NO ÂMBITO JURÍDICO EM PÂRAMETRO NACIONAL

2.2.1 Sistema Victor

Em 2018 o Supremo Tribunal Federal desenvolveu em parceria com a Universidade de Brasília o “Projeto Victor”, sistema de inteligência artificial projetado para auxiliar os ministros em seus julgamentos, cujo nome destinou-se a homenagear o ex-ministro da Corte, Victor Nunes Leal, por ser o primeiro ministro a tentar sistematizar os precedentes do STF.

Em sessão plenária do dia 30 de agosto de 2018, a Ministra Cármen Lúcia, então presidente do STF, anunciou que o sistema Victor já estava em funcionamento. A Ministra explicou que a máquina é utilizada para execução de quatro finalidades: I -

¹¹⁴ LEGAL SaaS AI Platform LawGeex levanta US \$ 7 milhões na rodada de financiamento. **In: CISION PRNewswire**, 2017. Disponível em: <https://www.prnewswire.com/news-releases/legal-saas-ai-platform-lawgeex-raises-7-million-in-funding-round-615570484.html>. Acesso em: 29 ago. 2021.

¹¹⁵ LAWGEEX. **Comparing the Performance of Artificial Intelligence to Human Lawyers in the Review of Standard Business Contracts**. 2018. Disponível em: <https://images.law.com/contrib/content/uploads/documents/397/5408/lawgeex.pdf>. Acesso em: 29 ago. 2021, p. 14.

conversão de imagens em textos no processo digital; II - separação do começo ao fim de um único documento em todo o acervo do tribunal; III - separação e classificação das peças processuais mais utilizadas nas atividades do STF; e, IV - a identificação dos temas de repercussão geral de maior incidência.¹¹⁶

A partir desse processamento, modelos de NLP (Natural Language Processing) são aplicados aos dados visando determinar em qual repercussão geral tal processo se encaixa. Houve a produção também de dois subprodutos ao projeto que são relevantes ao tribunal: transformação de imagens em textos para posteriores buscas e edições e outro classificador capaz de determinar automaticamente se uma peça jurídica é Recurso Extraordinário, Agravo em Recurso Extraordinário, Sentença, Acórdão, Despacho ou outra categoria genérica de documentos. Espera-se que uma vez em execução, o Victor contribua na celeridade e qualidade do fluxo de análises de processos jurídicos, sendo uma solução adequada às necessidades dos servidores e operadores do Direito do Supremo Tribunal Federal.¹¹⁷

Segundo Silva¹¹⁸, considerando a imensa quantidade de dados não estruturados, entre textos e imagens de documentos, com cerca de 350 novos processos de aproximadamente 60 páginas cada entrando todos os dias no STF, analisar e reconhecer padrões é realmente um grande desafio. Assim, num primeiro momento, Victor está trabalhando com cerca de 14.000 processos públicos e sem nenhum conteúdo sigiloso, entre aproximadamente 200 mil processos históricos. A escolha de atacar o problema através do *machine learning* foi absolutamente promissora, uma vez que os resultados obtidos com a tecnologia permitiram a realização das tarefas de classificações de textos com níveis de assertividade superiores a 93%. Esse avanço na separação de peças representa um futuro brilhante, que apoiará consideravelmente o trabalho diário de análise de processos dentro do STF. Toda a equipe envolvida, como professores e pesquisadores em geral, está bastante motivada em atuar no projeto devido à sua complexidade e ao potencial

¹¹⁶ AMORIM, Fernanda Pacheco. **Respeita as Mina:** Inteligência artificial e violência contra a mulher. 1. ed., Florianópolis: EMais Editora, 2019, p. 101.

¹¹⁷ INAZAWA, Pedro; PEIXOTO, Fabiano Hartmann; CAMPOS, Teófilo de; SILVA, Nilton; BRAZ, Fabricio. **Projeto Victor:** como o uso do aprendizado de máquina pode auxiliar a mais alta corte brasileira a aumentar a eficiência e a velocidade de avaliação judicial dos processos julgados. 2019. Disponível em: https://cic.unb.br/~teodecampos/ViP/inazawa_etal_compBrasil2019.pdf. Acesso em: 02 ago. 2021.

¹¹⁸ SILVA, Nilton Correia da. Notas iniciais sobre a evolução dos algoritmos do VICTOR: o primeiro projeto de inteligência artificial em supremas cortes do mundo. **In:** FERNANDES, Ricardo Vieira de Carvalho; CARVALHO, Angelo Gamba Prata de (Coord.). **Tecnologia jurídica & direito digital:** II Congresso Internacional de Direito, Governo e Tecnologia, p. 89-94, Belo Horizonte: Fórum, 2018.

produtivo que ele pode representar para o Tribunal e para o país.

O ministro Dias Toffoli, destacou que o sistema de inteligência artificial tem acuidade de 85%, sendo,

[...] responsável pela identificação de processos de repercussão geral é um mecanismo que converte imagens em texto, o que melhora e dinamiza a avaliação dos processos. Segundo o ministro Dias Toffoli, além de poupar tempo para o trabalho da Justiça, a nova ferramenta pode economizar recursos humanos. “O trabalho que custaria ao servidor de um tribunal entre 40 minutos e uma hora para fazer, o software faz em cinco segundos. Nossa ideia é replicar junto aos Tribunais Regionais Federais (TRFs), aos Tribunais de Justiça, aos Tribunais Regionais do Trabalho, enfim, trata-se de uma ferramenta para toda a magistratura”.¹¹⁹

O objetivo do projeto não é que o algoritmo tome a decisão final acerca da repercussão geral, mas sim que, com as máquinas “treinadas” para atuar em camadas de organização dos processos, os responsáveis pela análise dos recursos possam identificar os temas relacionados de forma mais clara e consistente.¹²⁰ Isso vai gerar, em consequência, mais qualidade e velocidade ao trabalho de avaliação judicial, com a redução das tarefas de classificação, organização e digitalização de processos.

Desde que foi instituída, pela Emenda Constitucional 45, de 2004, o instituto já registra 1.015 propostas de leading cases (casos paradigma) de repercussão geral. Desses casos, 682 tiveram a repercussão geral reconhecida e 325, negada. Até hoje, 379 causas de repercussão geral foram julgadas pelo Plenário do Supremo, 16 delas este ano. Outros 303 temas de repercussão geral aguardam julgamento pelo STF.¹²¹

Com a experiência do Projeto Victor surgem algumas perspectivas de que a IA e a tecnologia podem gerar, quando aplicadas ao Poder Judiciário. Dentre essas perspectivas do que pode ocorrer, tendo em conta as pesquisas que estão em curso, é de se ressaltar: a) a redução no tempo de tramitação de processos, em virtude da automação de procedimentos técnicos, o que fortalece, inclusive, a concretização do princípio da eficiência administrativa; b) o desenvolvimento de tecnologias e pesquisas genuinamente brasileiras, que levem em conta as particularidades do nosso congestionado sistema judicial; c) o incremento da agilidade e eficácia das ferramentas de consulta processual e jurisprudencial, o que gera também economia

¹¹⁹ CONSELHO NACIONAL DE JUSTIÇA. **Inteligência artificial:** Trabalho judicial de 40 minutos pode ser feito em 5 segundos. 2018. Disponível em: <https://www.cnj.jus.br/inteligencia-artificial-trabalho-judicial-de-40-minutos-pode-ser-feito-em-5-segundos/>. Acesso em: 28 ago. 2021, p. 01.

¹²⁰ SUPREMO TRIBUNAL FEDERAL. **Inteligência artificial vai agilizar a tramitação de processos no STF.** 2018. Disponível em: <http://www.stf.jus.br/portal/cms/verNoticiaDetalhe.asp?idConteudo=380038>. Acesso em: 10 ago. 2021.

¹²¹ CONSELHO NACIONAL DE JUSTIÇA. **Inteligência artificial:** Trabalho judicial de 40 minutos pode ser feito em 5 segundos. 2018. Disponível em: <https://www.cnj.jus.br/inteligencia-artificial-trabalho-judicial-de-40-minutos-pode-ser-feito-em-5-segundos/>. Acesso em: 28 ago. 2021, p. 01.

de tempo, precisão e coerência institucional; d) o tratamento isonômico das questões apresentadas ao Judiciário, que torna mais eficazes os princípios do contraditório, da ampla defesa e do livre acesso à justiça.¹²²

2.2.2 Projeto Sócrates

No mesmo sentido que o projeto Victor, em busca por integralizar a tecnologia e os meios administrativos para tornar o poder judiciário mais inteligente levaram o ministro João Otávio de Noronha, do Superior Tribunal de Justiça, a criar a Assessoria de Inteligência Artificial, o Sócrates, visando desenvolver melhorias no trabalho do STJ, principalmente quanto a gestão dos processos e seus volumes.¹²³

Um dos principais resultados dessa iniciativa é o Projeto Sócrates, construído com as ferramentas de IA. O Sócrates 1.0 – iniciado em maio de 2019 e já em operação em 21 gabinetes de ministros – faz a análise semântica das peças processuais com o objetivo de facilitar a triagem de processos, identificando casos com matérias semelhantes e pesquisando julgamentos do tribunal que possam servir como precedente para o processo em exame.¹²⁴

No relatório apresentado pela Fundação Getúlio Vargas em parceria com o Superior Tribunal de Justiça, consta como resultados do sistema Sócrates:

Redução do esforço na triagem de processos; apoio das atividades de análise de processos; e auxílio da seleção de representativos da controvérsia pelo Gabinete. É possível, fornecendo um caso-exemplo, identificar os demais processos que tratam da mesma matéria em um universo de 2 milhões de processos e 8 milhões de peças processuais, o que abrange todos os processos em tramitação no STJ e mais 4 anos de histórico, em 24 segundos. Além disso, é possível monitorar automaticamente os 1,5 mil novos processos que chegam diariamente ao Tribunal para seleção de matérias de interesse. Entre os ganhos já observados estão mais agilidade no julgamento, maior eficiência na seleção de precedentes qualificados e automatização da identificação de processos repetitivos que chegam ao Tribunal para

¹²² TOLEDO, Eduardo S. Projetos de inovação tecnológica na Administração Pública. *In*: FERNANDES, Ricardo Vieira de Carvalho; CARVALHO, Angelo Gamba Prata de (Coord.). **Tecnologia jurídica & direito digital**: II Congresso Internacional de Direito, Governo e Tecnologia. 2018. Belo Horizonte: Fórum, 2018.

¹²³ SUPERIOR TRIBUNAL DE JUSTIÇA. **Revolução tecnológica e desafios da pandemia marcaram gestão do ministro Noronha na presidência do STJ**. 2020. Disponível em: <https://www.stj.jus.br/sites/portalp/Paginas/Comunicacao/Noticias/23082020-Revolucao-tecnologica-e-desafios-da-pandemia-marcaram-gestao-do-ministro-Noronha-na-presidencia-do-STJ.aspx>. Acesso em: 04 ago. 2021.

¹²⁴ SUPERIOR TRIBUNAL DE JUSTIÇA. **Revolução tecnológica e desafios da pandemia marcaram gestão do ministro Noronha na presidência do STJ**. 2020. Disponível em: <https://www.stj.jus.br/sites/portalp/Paginas/Comunicacao/Noticias/23082020-Revolucao-tecnologica-e-desafios-da-pandemia-marcaram-gestao-do-ministro-Noronha-na-presidencia-do-STJ.aspx>. Acesso em: 04 ago. 2021.

juízo mais célere.¹²⁵

Posteriormente, o sistema foi aprimorado, visando trazer soluções para um dos principais desafios dos gabinetes, a identificação antecipada das controvérsias jurídicas do recurso especial, o Núcleo de Admissibilidade e Recursos Repetitivos, em conjunto com a equipe de IA do tribunal, projetou uma nova solução tecnológica, o Sócrates 2.0. A nova versão do Sócrates é capaz de apontar, de forma automática, o permissivo constitucional invocado para a interposição do recurso, os dispositivos de lei descritos como violados ou objeto de divergência jurisprudencial e os paradigmas citados para justificar a divergência.¹²⁶

De acordo com o relatório de pesquisa da FGV, as principais finalidades do Sócrates 2.0 são:

- identificação das controvérsias idênticas ou com abrangência delimitada para análise e afetação à sistemática dos recursos repetitivos;
- fomento de novas formas de triagem para potencializar o julgamento de mais processos em menos tempo, seja pelo impacto no Gabinete, nas Turmas ou nas Seções respectivas, bem como na Corte Especial;
- identificação dos casos com potencial de inadmissão para registro à Presidência;
- subsídio à Escola Corporativa do STJ nas definições de capacitação que melhor atendam à compreensão das matérias pendentes de julgamento.¹²⁷

Com a utilização da segunda versão do sistema, utilizando as linhas e linguagens de programação Python¹²⁸, Elastic¹²⁹ e Gensin¹³⁰, é possível relacionar

¹²⁵ SALOMÃO, Luis Felipe (coord.). **Inteligência Artificial: Tecnologia aplicada à gestão dos conflitos no âmbito do Poder Judiciário brasileiro**. Fundação Getúlio Vargas, 2019. Disponível em: https://ciapj.fgv.br/sites/ciapj.fgv.br/files/estudos_e_pesquisas_ia_1afase.pdf. Acesso em: 04 ago. 2021, p. 28.

¹²⁶ SUPERIOR TRIBUNAL DE JUSTIÇA. **Revolução tecnológica e desafios da pandemia marcaram gestão do ministro Noronha na presidência do STJ**. 2020. Disponível em: <https://www.stj.jus.br/sites/portalp/Paginas/Comunicacao/Noticias/23082020-Revolucao-tecnologica-e-desafios-da-pandemia-marcaram-gestao-do-ministro-Noronha-na-presidencia-do-STJ.aspx>. Acesso em: 04 ago. 2021.

¹²⁷ SALOMÃO, Luis Felipe (coord.). **Inteligência Artificial: Tecnologia aplicada à gestão dos conflitos no âmbito do Poder Judiciário brasileiro**. Fundação Getúlio Vargas, 2019. Disponível em: https://ciapj.fgv.br/sites/ciapj.fgv.br/files/estudos_e_pesquisas_ia_1afase.pdf. Acesso em: 04 ago. 2021, p. 28-29.

¹²⁸ Python é uma linguagem de programação criada por Guido van Rossum em 1991. Os objetivos do projeto da linguagem eram: produtividade e legibilidade. Em outras palavras, Python é uma linguagem que foi criada para produzir código bom e fácil de manter de maneira rápida. Entre as características da linguagem que ressaltam esses objetivos estão: baixo uso de caracteres especiais, o que torna a linguagem muito parecida com pseudo-código executável; o uso de indentação para marcar blocos; quase nenhum uso de palavras-chave voltadas para a compilação; coletor de lixo para gerenciar automaticamente o uso da memória; etc.

¹²⁹ O Elasticsearch é um mecanismo de busca e análise de dados distribuído e open source para todos os tipos de dados, incluindo textuais, numéricos, geoespaciais, estruturados e não estruturados. O Elasticsearch é desenvolvido sobre o Apache Lucene e foi lançado pela primeira vez em 2010 pela Elasticsearch N.V.

¹³⁰ É um projeto especificamente para lidar com grandes coleções de texto digitais brutos e não estruturados, usando streaming de dados e algoritmos incrementais eficientes, como Análise

processos que se assemelhem ao processo em pauta com eficácia acima de 95 por cento de similaridade, conforme demonstra Luiz Anísio Vieira Batitucci e Amilar Domingos Moreira Martins.¹³¹

2.2.3 Plataforma de Inteligência Artificial Dra. Luzia

A Dra. Luzia é a primeira Robô Advogada Assistente Brasileira, criada pela *startup* brasileira Legal Labs, trata-se de um sistema de inteligência artificial projetado para acelerar a tramitação de processos de execução fiscal na Procuradoria Geral do Distrito Federal (PGDF).¹³²

O objetivo é que a Dra. Luzia atenda procuradorias de Estado, e ajude o órgão a fazer peticionamentos automático a partir de *machine learning*, como também gestão de processos jurídicos e acompanhamento de resultados. "Em uma semana, a Dra. Luzia gerou 668 petições de um total de 773. Se fizer uma regra de três, a Robô Advogada Assistente fez de 85% a 90% de todo o trabalho. Toda realização dessas tarefas foi feita sem nenhum contato humano", informou o procurador.¹³³

No entanto, é importante mencionar que a Robô Advogada se enquadra na IA Fraca, pois não possui capacidade cognitiva, desenvolvendo suas atividades mediante sua programação e cruzamento de dados.

Dra. Luzia foi desenvolvida com as seguintes capacidades:

- 1 - entender os processos, o seu andamento e quais suas possíveis soluções,
- 2 - cruzar dados e encontrar endereços ou bens dos envolvidos nos processos,
- 3 - auxiliar na tramitação dos processos de execuções fiscais,
- 4 - avaliar aproximadamente 300 mil processos em andamento de cobrança

Semântica Latente, Alocação de Dirichlet Latente e Projeções Aleatórias, descoberta da estrutura semântica dos documentos examinando padrões estáticos de ocorrência da palavras dentro de um corpus de documentos de treinamento. Esses algoritmos não são supervisionados, o que significa que nenhuma entrada humana é necessária, necessitando apenas de um corpus de documentos de texto simples, o que o diferencia da maioria dos outros pacotes de software científicos que visam somente o processamento em lote de memória.

¹³¹ BATITUCCI, Luiz Anísio Vieira; MARTINS, Amilar Domingos Moreira. **Sistemas apoiados por Inteligência Artificial**: Superior Tribunal de Justiça. 2019. Disponível em: <http://www.jf.gov.br/cjf/corregedoria-da-justica-federal/centro-de-estudos-judiciarios-1/eventos/eventos-cej/2019/stj-apresentacao-enatic-junho-2019-2.pdf>. Acesso em: 04 ago. 2021.

¹³² COELHO, João Victor de Assis Brasil Ribeiro. **Aplicações e Implicações da Inteligência Artificial no Direito**. 2017. 61 f. TCC (Graduação em Direito) – Universidade de Brasília, Brasília, 2017. Disponível em: https://bdm.unb.br/bitstream/10483/18844/1/2017_JoaoVictordeAssisBrasilRibeiroCoelho.pdf. Acesso em: 29 ago. 2021.

¹³³ COSTA, Suzana Rita da. **A Contribuição da Inteligência Artificial na Celeridade dos Trabalhos Repetitivos no Sistema Jurídico**. 2020. 72 fls. Dissertação (Mestrado em Mídia e Tecnologia) – Universidade Estadual Paulista Júlio de Mesquita Filho – Unesp, Bauru, 2020, p. 40.

de dívida ativa do órgão,
 5 - interpretar decisões dos juízes no despacho e certidões,
 6 - definir qual a petição mais adequada para cada situação jurídica apresentada de acordo com sua programação,
 7 - gerar novas petições a partir destas.¹³⁴

Segundo Ricardo Fernandes, fundador da Legallabs e desenvolvedor do sistema, os resultados obtidos indicaram que a Dra. Luzia é capaz de processar 1000 petições em 1 minuto e 56 segundos, em média, ao passo que 4 servidores levavam 4 dias úteis para realizar a mesma tarefa.¹³⁵

Nesse sentido, importante destacar que as execuções fiscais representam 39% do total de casos pendentes, que por sua vez, correspondem a 74% de execuções pendentes em todo Poder Judiciário, com a maior taxa de congestionamento na prestação jurisdicional, de 91,7%. Exemplificando: a cada 100 processos de execução fiscal, somente 8 foram baixados. Assim, o “[...] tempo de giro do acervo desses processos é de 11 anos, ou seja, mesmo que o Judiciário parasse de receber novas execuções fiscais, ainda seriam necessários 11 anos para liquidar o acervo existente”.¹³⁶

De acordo com Ricardo Fernandes o procurador do Distrito Federal, idealizador e também usuário de Dra. Luzia, com a inserção da IA no setor os processos têm sido realizados em uma velocidade que o ser humano não consegue alcançar. E assim, uma diminuição grande no acúmulo de trabalho deste setor. Além disso, houve uma diminuição na margem de erros. Com a IA, os trabalhos se tornaram mais céleres e assertivos, e assim, decisões muito mais rápidas.¹³⁷

2.2.4 Sistemas Alice, Sofia e Mônica

Alice, Sofia e Monica são três robôs que auxiliam o trabalho do Tribunal de Contas da União na tarefa de fiscalizar a existência de fraudes em licitações. Além da

¹³⁴ COSTA, Suzana Rita da. **A Contribuição da Inteligência Artificial na Celeridade dos Trabalhos Repetitivos no Sistema Jurídico**. 2020. 72 fls. Dissertação (Mestrado em Mídia e Tecnologia) – Universidade Estadual Paulista Júlio de Mesquita Filho – Unesp, Bauru, 2020, p. 40.

¹³⁵ SERAPIÃO, Fábio. Dra. Luzia. *In: Estadão*, 13 de maio de 2018. Disponível em: <https://politica.estadao.com.br/blogs/fausto-macedo/dra-luzia/>. Acesso em: 29 ago. 2021.

¹³⁶ CONSELHO NACIONAL DE JUSTIÇA. **Relatório Justiça em Números 2018: ano-base 2017**. 2018. Disponível em: www.cnj.jus.br/files/conteudo/arquivo/2018/09/8d9faee7812d35a58cee3d92d2df2f25.pdf. Acesso em: 29 ago. 2021, p. 125.

¹³⁷ FERNANDES, R. Dr^a Luzia, primeira robô-advogada do Brasil, já tem trabalho pela frente. *In: Instituto Humanista Unisinos*, 2017. Disponível em: <http://www.ihu.unisinos.br/78-noticias/569427-dr-luzia-primeira-robo-advogada-do-brasil-jatem-trabalho-pela-frente>. Acesso em: 28 ago. 2021.

análise dos editais, elas ajudam os auditores na hora de redigir textos e ficam de olho nas contratações públicas. O Tribunal de Contas da União tem se utilizado de sistemas inteligentes, no âmbito operacional interno, para aumentar sua produtividade. A robô Alice, abreviação para Análise de Licitações e Editais, com o auxílio de outros dois robôs, Sofia e Monica, faz uma varredura nas contratações federais, a fim de detectar possíveis irregularidades.¹³⁸

Alice, Sofia e Monica são interfaces de um sistema maior, o Laboratório de Informações de Controle (Labcontas), que reúne 77 (setenta e sete) bases de dados integradas entre si, das quais são exemplo, registro de compras governamentais, lista de políticas públicas, composição societária de empresas, entre outros. O sistema Labcontas possui, ainda, tecnologia que permite o cruzamento de informações entre as bases de dados que o compõem. Além da Corte de Contas da União, a Controladoria Geral da União, o Ministério Público Federal, a Polícia Federal e Tribunais de Contas estaduais, também utilizam Alice, Sofia e Monica para o desempenho de tarefas.¹³⁹

2.2.4.1 Alice

No ar desde fevereiro de 2017, Alice lê editais de licitações e atas de registros de preços publicados pela administração federal, além de alguns órgãos públicos estaduais e empresas estatais, através da coleta de informações no Diário Oficial e no Comprasnet¹⁴⁰. A partir dessa varredura, Alice emite um relatório indicando ao auditor indícios de irregularidades, a fim de que ele possa analisar o edital ou a ata de forma mais detalhada. Com a ajuda da Alice, os auditores conseguiram suspender contratações irregulares em Estados e até em editais do Itamaraty, demonstrando a contribuição do sistema computacional para a otimização, agilidade e eficiência do

¹³⁸ GOMES, Helton Simões. Como as robôs Alice, Sofia e Monica ajudam o TCU a caçar irregularidades em licitações. *In: G1*, 18 mar. 2018. Disponível em: <https://g1.globo.com/economia/tecnologia/noticia/como-as-robos-alice-sofia-e-monica-ajudam-o-tcu-a-cacar-irregularidades-em-licitacoes.ghtml>. Acesso em: 29 ago. 2021.

¹³⁹ GOMES, Helton Simões. Como as robôs Alice, Sofia e Monica ajudam o TCU a caçar irregularidades em licitações. *In: G1*, 18 mar. 2018. Disponível em: <https://g1.globo.com/economia/tecnologia/noticia/como-as-robos-alice-sofia-e-monica-ajudam-o-tcu-a-cacar-irregularidades-em-licitacoes.ghtml>. Acesso em: 29 ago. 2021.

¹⁴⁰ Comprasnet é o Portal de Compras do Governo Federal, instituído com o objetivo de disponibilizar para a sociedade informações referentes às licitações e contratações realizadas pelo governo federal, bem como para permitir a realização de processos eletrônicos de aquisição.

serviço público prestado pelo órgão.¹⁴¹

O sistema Alice surgiu para resolver alguns problemas, no que se refere as Licitações:

- 1) Inúmeras quantidades de órgãos que realizam licitações;
- 2) Número de licitações por ano desde 2000. Mais de 60 mil licitações por ano na União, dá mais de 200 licitações diárias;
- 3) Crescimento da utilização dos pregões eletrônicos – dificuldade maior. Publicação: dia zero – em 15 dias será realizada a sessão. Então em 25 dias pode ser que já tenha sido homologada a licitação.
- 4) A fase pública da licitação (após a publicação do edital) é a fase em que o auditor pode tomar alguma medida de fiscalização eficiente, já que é anterior à assinatura do contrato administrativo.
- 5) Na União existe a base de dados do SIASG que tem uma defasagem de 30 dias da publicação do edital.
- 6) A informação mais atualizada que se consegue encontrar é no comprasnet.
- 7) O ALICE veio resolver esses problemas, isto é, ESCOLHER MELHOR O QUE FAZER – ou seja, busca identificar automaticamente indícios de irregularidades nas licitações de forma tempestiva, tão logo ocorra sua publicação no portal do comprasnet.¹⁴²

Wesley Vaz Silva, diretor da Secretaria de Fiscalização de Tecnologia da Informação do TCU, descreve como funciona a atuação de Alice ao dizer que é um robô, usado pelo TCU para caçar fraudes e outras irregularidades em licitações. A Alice faz as análises das licitações e editais, coleta de informações dos Diários Oficiais, lê editais de licitações e atas de registro de preços publicados pela administração federal, além de alguns órgãos públicos estaduais e estaduais, coleta informações do Diário Oficial e do Comprasnet (o sistema que registra as compras governamentais), a partir daí, elabora um documento e aponta aos auditores se há indícios de desvios, os auditores do Tribunal de Contas da União recebem pontualmente às 19h, todos os dias, um e-mail de Alice com os informes das centenas de contratações federais publicadas. Leva as informações aos membros do TCU e o número de processos por estado, apontando a margem de erro e o valor de cada um, e nas informações ela apresenta quais contratações podem conter irregularidades e se há indícios de fraude.¹⁴³

¹⁴¹ DESORDI, Danubia; BONA, Carla Della. A inteligência artificial e a eficiência na administração pública. **In: Revista de Direito**, Viçosa, vol. 02, n. 02, 2020. Disponível em: <https://periodicos.ufv.br/revistadir/article/view/9112/5928>. Acesso em: 29 ago. 2021.

¹⁴² NAVES, Fernanda de Moura Ribeiro. **Relatório de visita realizada ao Tribunal de Contas da União para exposição sobre os Sistemas Labcontas, Alice, Sofia e Mônica**. Goiânia, 2018. Disponível em: https://files.cercomp.ufg.br/weby/up/949/o/Relat%C3%B3rio_de_visita_ao_TCU_sobre_ALICE.pdf. Acesso em: 29 ago. 2021, p. 07.

¹⁴³ GOMES, Helton Simões. Como as robôs Alice, Sofia e Monica ajudam o TCU a caçar irregularidades em licitações. **In: G1**, 18 mar. 2018. Disponível em: <https://g1.globo.com/economia/tecnologia/noticia/como-as-robos-alice-sofia-e-monica-ajudam-o-tcu-a-caçar-irregularidades-em-licitacoes.ghtml>. Acesso em: 29 ago. 2021.

Um fato muito interessante é que racionaliza o trabalho dos auditores do TCU é que dentro de cada tipologia (ou irregularidade encontrada), o Alice já fornece a fundamentação legal e jurisprudencial pertinente para a formulação de representação das irregularidades perante aquele Tribunal, caso não seja possível sanar tais irregularidades via telefone ou e-mail.¹⁴⁴

O TCU internamente considera que o Alice é uma ferramenta em constante desenvolvimento. Algumas das melhorias previstas são o controle de dispensa e inexigibilidade (que será feito por meio do DOU), RDC eletrônico, que já se encontram em desenvolvimento, assim como, a construção de novas tipologias e o acréscimo de novas fontes de informação sobre licitações, além do Comprasnet.¹⁴⁵

2.2.4.2 Sofia

O Sofia (Sistema de Orientação sobre Fatos e Indícios para o Auditor) é uma ferramenta que provê informações ao auditor no momento da elaboração de documentos de controle externo. Por meio desse sistema é feita revisão nos relatórios de auditoria e instruções em geral, além de ser efetuada busca de correlação das informações neles constantes. O Sofia, por exemplo, capta as informações associadas aos CNPJs indicados no documento e verifica se já foram aplicadas sanções àquelas empresas ou se elas já foram responsabilizadas em outros processos em trâmite no TCU, ou, ainda, elenca os contratos já pactuados por essas empresas com órgãos ou entidades da Administração Pública Federal, entre outras informações providas.¹⁴⁶

De acordo com Suzana Rita da Costa, Sofia realiza as seguintes atividades:

- 1 - realiza um trabalho mais assertivo, apontando supostos erros nos textos dos auditores, indicando outras fontes de referências,
- 2 - sugere a respeito de informações relacionadas as partes envolvidas

¹⁴⁴ NAVES, Fernanda de Moura Ribeiro. **Relatório de visita realizada ao Tribunal de Contas da União para exposição sobre os Sistemas Labcontas, Alice, Sofia e Mônica.** Goiânia, 2018. Disponível em: https://files.cercomp.ufg.br/weby/up/949/o/Relat%C3%B3rio_de_visita_ao_TCU_sobre_ALICE.pdf. Acesso em: 29 ago. 2021, p. 09.

¹⁴⁵ NAVES, Fernanda de Moura Ribeiro. **Relatório de visita realizada ao Tribunal de Contas da União para exposição sobre os Sistemas Labcontas, Alice, Sofia e Mônica.** Goiânia, 2018. Disponível em: https://files.cercomp.ufg.br/weby/up/949/o/Relat%C3%B3rio_de_visita_ao_TCU_sobre_ALICE.pdf. Acesso em: 29 ago. 2021, p. 09.

¹⁴⁶ COSTA, Marcos Bemquerer; BASTOS, Patrícia Reis Leitão. Alice, Monica, Adele, Sofia, Carina e Ágata: o uso da inteligência artificial pelo Tribunal de Contas da União. **In: Revista do Tribunal de Contas do Estado de Goiás**, Belo Horizonte, ano 2, n. 3, p. 11-34, jan./jun. 2020.

ao tema em questão, por exemplo, se um texto propõe punição, ela pode indicar se há sanções contra a companhia, se ela tem processo no tribunal, e ainda apontar se a empresa possui outros contratos com a administração pública,

3 - e também cria alerta com os dados como por exemplo de um CPF registrado pelo auditor, a existência e a validade de contratos de uma entidade, se há registro de óbito sobre determinada pessoa, e se o cidadão ou empresa está ou não cadastrado no sistema do TCU. "A Sofia é um botãozinho no Word", explica a auditora federal Tania Lopes Pimenta Cioato. Ao apertá-lo, ela lista informações associadas aos números de CNPJ, do processo e de CPF incluídos no texto. Diz até se algum dos envolvidos já morreu.¹⁴⁷

Sendo assim, esse sistema efetua uma análise crítica dos textos produzidos pelos auditores (relatórios, instruções, pareceres, entre outros), identificando os principais elementos e confrontando-os com as informações que estão nos sistemas à disposição do TCU. Além disso, o Sofia permite encontrar correlações relevantes e indícios de erros e irregularidades que, caso não corrigidos a tempo, poderão comprometer a confiabilidade dos documentos gerados pelas unidades técnicas da Corte de Contas.¹⁴⁸

Entende-se que o Sofia representa mais que uma solução, a ideia pode ser aplicada a todos os casos em que peças jurídicas precisam ser construídas e que possam ter sua análise de mérito mais efetiva a partir do consumo de informações sobre os componentes dos textos, tais como pessoas físicas e jurídicas.

2.2.4.3 Monica

Monica consiste em um painel onde todas as compras públicas podem ser visualizadas, incluindo as que são ignoradas por Alice, como contratações diretas e inexigibilidades de licitação.¹⁴⁹

Nesse sentido, Suzana Rita da Costa colabora salientando que, conforme já mencionado, a Robô Mônica é um tipo de painel que mostra todas as compras públicas realizadas, inclusive algumas que Alice não tenha percebido, como por exemplo:

¹⁴⁷ COSTA, Suzana Rita da. **A Contribuição da Inteligência Artificial na Celeridade dos Trabalhos Repetitivos no Sistema Jurídico**. 2020. 72 fls. Dissertação (Mestrado em Mídia e Tecnologia) – Universidade Estadual Paulista Júlio de Mesquita Filho – Unesp, Bauru, 2020, p. 48-49.

¹⁴⁸ COSTA, Marcos Bemquerer; BASTOS, Patrícia Reis Leitão. Alice, Monica, Adele, Sofia, Carina e Ágata: o uso da inteligência artificial pelo Tribunal de Contas da União. **In: Revista do Tribunal de Contas do Estado de Goiás**, Belo Horizonte, ano 2, n. 3, p. 11-34, jan./jun. 2020.

¹⁴⁹ DESORDI, Danubia; BONA, Carla Della. A inteligência artificial e a eficiência na administração pública. **In: Revista de Direito**, Viçosa, vol. 02, n. 02, 2020. Disponível em: <https://periodicos.ufv.br/revistadir/article/view/9112/5928>. Acesso em: 29 ago. 2021.

- 1 - contratações diretas, e aquelas que são feitas por meio de inexigibilidade de licitação, que é quando um certo tipo de serviço ou produto só existe um fornecedor,
- 2 - traz informações sobre as compras públicas na esfera Federal, incluindo os poderes Executivo, Legislativo e Judiciário, além do Ministério Público,
- 3 - Mônica também faz um trabalho mensal de obtenção de dados, de uma forma muito rápida por meio de palavras-chave no objeto das aquisições. Tudo graças a tecnologia inteligente.¹⁵⁰

O Monica (Monitoramento Integrado para o Controle de Aquisições) é um painel que contempla informações relativas às aquisições efetuadas pela esfera federal, incluindo os poderes Executivo, Legislativo e Judiciário, além do Ministério Público Federal. Por enquanto, esse monitoramento se restringe às aquisições efetuadas no âmbito do Sistema Integrado de Administração de Serviços Gerais (Siasg), de tal forma que ainda não estão incluídas nesse painel as compras realizadas por empresas estatais e aquelas efetivadas por meio do Regime Diferenciado de Contratações – RDC.¹⁵¹

O robô faz um trabalho mensal de obtenção de dados, com exceção das informações sobre pregões que são atualizadas semanalmente. Além disso, a tecnologia permite que sejam feitas buscas rápidas por palavras-chave no objeto das aquisições. Contudo, importante mencionar que apensar de ser chamada constantemente como robô, diferentemente de Alice, Monica não é um robô, ainda é apenas um software.

2.3 REGULAMENTAÇÃO DA INTELIGÊNCIA ARTIFICIAL NO MUNDO E NO BRASIL

O mercado de inteligência artificial está em constante evolução e cada vez mais acelerado, e, com isso, a discussão sobre o arcabouço regulatório da tecnologia também ganha destaque. Isto porque, conforme mais pessoas e negócios adotam sistemas inteligentes, torna-se inevitável a determinação de normas para que seu uso seja responsável e seguro.

O uso da inteligência artificial no judiciário não é diferente, visto que apesar de ser evidente que sua implementação é de suma importância para a garantia de

¹⁵⁰ COSTA, Suzana Rita da. **A Contribuição da Inteligência Artificial na Celeridade dos Trabalhos Repetitivos no Sistema Jurídico**. 2020. 72 fls. Dissertação (Mestrado em Mídia e Tecnologia) – Universidade Estadual Paulista Júlio de Mesquita Filho – Unesp, Bauru, 2020, p. 49.

¹⁵¹ COSTA, Marcos Bemquerer; BASTOS, Patrícia Reis Leitão. Alice, Monica, Adele, Sofia, Carina e Ágata: o uso da inteligência artificial pelo Tribunal de Contas da União. **In: Revista do Tribunal de Contas do Estado de Goiás**, Belo Horizonte, ano 2, n. 3, p. 11-34, jan./jun. 2020.

uma maior eficiência, de tempo e de recurso, nos procedimentos dentro de um processo judicial até sua decisão definitiva, existe uma grande preocupação desde a redução dos postos de trabalho até a possibilidade de seus algoritmos criarem vieses involuntários, reproduzindo padrões contrários à legislação.

Atualmente, ainda há uma escassez de regulamentação sobre a implementação e uso da inteligência artificial, contudo, já se pode vislumbrar grandes passos nessa direção no Brasil e no mundo, como é o caso da União Europeia que vem se mostrando pioneira nas diretrizes para o uso da tecnologia de inteligência artificial.

2.3.1 As propostas para regulamentar a inteligência artificial ao redor do mundo

Primeiramente, é importante mencionar o que motivou a urgência em se definir uma legislação sobre a inteligência artificial no mundo. Em 2016, a Organização de Cooperação pelo Desenvolvimento Econômico (OCDE) já divulgava um documento no qual listava preocupações em relação ao uso da inteligência artificial. Eram elas: “a) Aumento do desemprego devido à automação; b) Maior desequilíbrio na distribuição de renda; c) Resultados enviesados pela ausência de supervisão humana”¹⁵². Diante disso, em maio de 2019, a OCDE criou um guia de recomendações para o desenvolvimento de inteligência artificial, no qual apresenta cinco princípios para o desenvolvimento de inteligência artificial, que são:

- a) A IA deve operar em benefício das pessoas e do planeta;
- b) Sistemas com IA devem respeitar os Direitos Humanos, valores democráticos e a diversidade;
- c) Sistemas de IA devem permitir intervenção humana quando necessário;
- d) Sistemas de IA devem ser transparentes, na medida do possível, para que cientistas entendam suas decisões;
- e) Riscos relacionados à IA devem ser constantemente avaliados.¹⁵³

Neste contexto diversos países tomaram a iniciativa de criar regulamentações sobre o uso da inteligência artificial, dentre os principais estão a China, o Brasil (que se abordará em tópico específico), Estados Unidos e União Europeia.

Sendo assim, logo após a OCDE criar o guia de recomendações para o

¹⁵² REGULAMENTAÇÃO da inteligência artificial: entenda o debate mundial. *In: Alana*, [s.d.]. Disponível em: <https://blog.alana.ai/cresca-com-alana/regulamentacao-da-inteligencia-artificial-entenda-o-debate-mundial/>. Acesso em: 22 out. 2021.

¹⁵³ REGULAMENTAÇÃO da inteligência artificial: entenda o debate mundial. *In: Alana*, [s.d.]. Disponível em: <https://blog.alana.ai/cresca-com-alana/regulamentacao-da-inteligencia-artificial-entenda-o-debate-mundial/>. Acesso em: 22 out. 2021.

desenvolvimento de inteligência artificial, em junho de 2019, a China divulgou regras de governança, lançando 8 princípios que devem ser seguidos no desenvolvimento de inteligência artificial no país. Esses princípios possuem o objetivo de promover o desenvolvimento saudável de uma nova geração de IA; coordenar melhor a relação entre desenvolvimento e governança; garantir que a IA seja segura/protegida, confiável e controlável; promover o desenvolvimento econômica, social e ecologicamente sustentável; e juntos construir uma comunidade de destino comum para a humanidade.¹⁵⁴

Os oito princípios de governança de IA propostos são:

- a) harmonia e simpatia;
- b) equidade e justiça;
- c) inclusão e compartilhamento;
- d) respeito pela privacidade;
- e) segurança e controlabilidade;
- f) responsabilidade compartilhada;
- g) cooperação aberta;
- h) governança ágil;¹⁵⁵

Basicamente, o documento trata sobre a necessidade de eliminar por meio de tecnologias e métodos de gerenciamento aprimorados as questões de Viés e discriminação no processo de aquisição de dados, projeto de algoritmo, desenvolvimento de tecnologia, bem como a pesquisa, o desenvolvimento e a aplicação do produto. Ainda, o desenvolvimento de IA deve respeitar e proteger a privacidade pessoal e proteger o direito dos indivíduos de saber e de escolha. Devem ser estabelecidos padrões para a coleta, armazenamento, processamento e uso de informações pessoais. Os mecanismos de revogação de autorização de informações pessoais devem ser melhorados. Por fim, o roubo, a adulteração, a divulgação ilegal e qualquer outra coleta ou uso ilegal de informações pessoais devem ser combatidos.¹⁵⁶

Já o governo dos Estados Unidos publicou em janeiro de 2020 uma proposta de regulamentação da inteligência artificial, a qual apresenta dez princípios para as agências governamentais aderirem ao propor regulamentos de Inteligência Artificial

¹⁵⁴ CHINA: Lançados Princípios de Governança de IA. *In: Library of Congress*, 2019. Disponível em: <https://www.loc.gov/item/global-legal-monitor/2019-09-09/china-ai-governance-principles-released/>. Acesso em: 22 out. 2021.

¹⁵⁵ CHINA: Lançados Princípios de Governança de IA. *In: Library of Congress*, 2019. Disponível em: <https://www.loc.gov/item/global-legal-monitor/2019-09-09/china-ai-governance-principles-released/>. Acesso em: 22 out. 2021.

¹⁵⁶ CHINA: Lançados Princípios de Governança de IA. *In: Library of Congress*, 2019. Disponível em: <https://www.loc.gov/item/global-legal-monitor/2019-09-09/china-ai-governance-principles-released/>. Acesso em: 22 out. 2021.

para o setor privado. Os princípios propostos pela OSTP têm três objetivos principais: a) garantir o envolvimento do público; b) limitar o alcance regulatório; e c) desenvolver uma Inteligência Artificial confiável, segura e transparente.¹⁵⁷

Os princípios, que foram redigidos de maneira ampla para orientar futuros regulamentos, são:

- a) confiança pública na IA;
- b) participação do público;
- c) integridade científica e qualidade de informação;
- d) avaliação e gerenciamento de riscos;
- e) custos e benefícios;
- f) flexibilidade;
- g) justiça e não discriminação;
- h) divulgação e transparência;
- i) segurança e proteção;
- j) coordenação interinstitucional.¹⁵⁸

Atualmente, a Casa Branca está estudando a “AI Bill of Rights”, (declaração de direitos da inteligência artificial). O objetivo é limitar os usos potencialmente prejudiciais dessa tecnologia. O nome faz referência a “Bill of Rights” –a declaração dos direitos fundamentais dos Estados Unidos, de 1789. Segundo o Axios, o departamento de Política Científica e Tecnológica da Casa Branca lançou no dia 08 outubro de 2021 um trabalho de apuração para examinar o reconhecimento facial e outras ferramentas tecnológicas.¹⁵⁹

Dentro do contexto europeu, pode-se mencionar a Carta de ética europeia sobre o uso da inteligência artificial em sistemas judiciais e seu ambiente, elaborada pela Comissão Europeia para Eficiência da Justiça e adotada na 31ª sessão plenária em Estrasburgo, nos dias 3 e 4 de dezembro de 2018.

A Comissão apresenta, por meio da Carta Europeia sobre ética no uso da inteligência artificial, uma estrutura de princípios que podem orientar os formuladores de políticas, legisladores e profissionais da justiça em relação ao rápido desenvolvimento da IA nos processos judiciais. O objetivo é garantir que tal tecnologia continue a serviço do interesse geral e que seu uso respeite os direitos individuais.

Nesse passo, os cinco princípios que o documento traz são:

¹⁵⁷ CASA Branca divulga princípios para a regulamentação de Inteligência Artificial. **In: Internetlab**, 2020. Disponível em: <https://www.internetlab.org.br/pt/itens-semanario/eua-casa-branca-divulga-principios-para-a-regulamentacao-de-inteligencia-artificial/>. Acesso em: 22 out. 2021.

¹⁵⁸ CASA Branca divulga princípios para a regulamentação de Inteligência Artificial. **In: Internetlab**, 2020. Disponível em: <https://www.internetlab.org.br/pt/itens-semanario/eua-casa-branca-divulga-principios-para-a-regulamentacao-de-inteligencia-artificial/>. Acesso em: 22 out. 2021.

¹⁵⁹ Casa Branca estuda ‘declaração de direitos da inteligência artificial’. In: Poder 360. 2021. Disponível em: <https://www.poder360.com.br/tecnologia/casa-branca-estuda-declaracao-de-direitos-da-inteligencia-artificial/>. Acesso em: 22 out. 2021.

1. Princípio do respeito aos direitos fundamentais: garantir que o design e implementação de ferramentas de inteligência artificial e os serviços são compatíveis com os direitos fundamentais.
2. Princípio da não discriminação: prevenir especificamente o desenvolvimento ou intensificação de qualquer discriminação entre indivíduos ou grupos de indivíduos.
3. Princípio da qualidade e segurança: no que diz respeito ao tratamento de decisões judiciais e de dados, utilizar fontes certificadas e dados intangíveis com modelos elaborados de forma multidisciplinar, em ambiente tecnológico seguro.
4. Princípio da transparência, imparcialidade e justiça: faça métodos de processamento de dados acessíveis e compreensíveis, autorizar auditorias externas.
5. Princípio "sob controle do usuário": impede uma prescrição abordagem e garantir que os usuários sejam atores informados e em controle das escolhas feitas. (Tradução nossa)¹⁶⁰

Salienta-se que o cumprimento desses princípios no processamento de decisões e dados judiciais por algoritmos se mostra necessário como forma de garantir o Estado Democrático de Direito.

Portanto, o objeto da Carta se refere especificamente ao processamento de decisões e dados judiciais por inteligência artificial. Por exemplo, as ferramentas que visam apoiar os profissionais do direito na realização de pesquisas jurídicas ou na antecipação do possível resultado de um caso apresentado a um tribunal (os chamados instrumentos de “justiça preditiva”). Também, aquelas desenvolvidas para apoiar as Cortes no gerenciamento de casos (por exemplo, verificando e atribuindo peças às seções responsáveis dos tribunais) ou para analisar o seu desempenho. Podem ser usadas fora do processo de litígio, na estrutura, por exemplo, de solução de disputas on-line.¹⁶¹

Com base na Carta de Ética, em 2019, foi elaborado pelo grupo de peritos de alto nível sobre a inteligência artificial (GPAN IA), instituição criada pela Comissão Europeia, um guia com as orientações éticas para uma IA de confiança, que prevê

¹⁶⁰ Texto original: 1 Principle of respect for fundamental rights: ensure that the design and implementation of artificial intelligence tools and services are compatible with fundamental rights. 2 Principle of non-discrimination: specifically prevent the development or intensification of any discrimination between individuals or groups of individuals. 3 Principle of quality and security: with regard to the processing of judicial decisions and data, use certified sources and intangible data with models elaborated in a multi-disciplinary manner, in a secure technological environment. 4 Principle of transparency, impartiality and fairness: make data processing methods accessible and understandable, authorise external audits. 5 Principle “under user control”: preclude a prescriptive approach and ensure that users are informed actors and in control of the choices made. (CONSELHO EUROPEU. European Commission for the Efficiency of Justice (CEPEJ). **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-december-2018/16808f699c>. Acesso em: 08 set. 2021).

¹⁶¹ LAGE, Fernanda de Carvalho. **Manual de Inteligência Artificial no Direito Brasileiro**. Salvador: Editora Juspodivm, 2021.

pilares como supervisão humana, segurança e robustez técnica, privacidade e governança de dados, transparência e não-discriminação na construção de sistemas inteligentes. Buscando evitar que a leitura fique cansativa e repetitiva, o referido documento será melhor estudado no terceiro capítulo.

Mais recentemente, em 21 de abril de 2021, a Comissão Europeia apresentou a Proposta de Regulamento do Parlamento Europeu e do Conselho que estabelece regras harmonizadas em matéria de Inteligência Artificial e altera determinados atos legislativos da União, trazendo novas regras para promover a confiança na inteligência artificial.¹⁶²

As novas regras serão aplicadas diretamente e da mesma forma em todos os Estados-Membros, com base numa definição de inteligência artificial orientada para o futuro. A presente Proposta regulamentar em matéria de inteligência artificial traz os seguintes objetivos específicos:

- garantir que os sistemas de IA colocados no mercado da União e utilizados sejam seguros e respeitem a legislação em vigor em matéria de direitos fundamentais e valores da União,
- garantir a segurança jurídica para facilitar os investimentos e a inovação no domínio da IA,
- melhorar a governação e a aplicação efetiva da legislação em vigor em matéria de direitos fundamentais e dos requisitos de segurança aplicáveis aos sistemas de IA,
- facilitar o desenvolvimento de um mercado único para as aplicações de IA legítimas, seguras e de confiança e evitar a fragmentação do mercado.¹⁶³

A proposta, segundo os legisladores europeus, pode ir além. Existe a possibilidade de criação de uma agência reguladora de IA e até a criação de um modelo de gerenciamento de risco na IA com base em quatro critérios: inaceitável, elevado, limitado e mínimo. Nessas definições, seriam enquadradas práticas de risco como brinquedos que podem atuar como manipuladores do comportamento de crianças, sistemas de ranking personalizados com base em características pessoais, envios de spam ou qualquer outro artifício que possa ferir liberdades individuais.¹⁶⁴

¹⁶² COMISSÃO EUROPEIA. **Proposta de regulamento que estabelece regras harmonizadas em matéria de Inteligência Artificial (Regulamento Inteligência Artificial) e altera determinados atos legislativos da União.** 21 abr. 2021. Disponível em: https://eur-lex.europa.eu/resource.html?uri=cellar:e0649735-a372-11eb-9585-01aa75ed71a1.0004.02/DOC_1&format=PDF. Acesso em: 08 set. 2021.

¹⁶³ COMISSÃO EUROPEIA. **Proposta de regulamento que estabelece regras harmonizadas em matéria de Inteligência Artificial (Regulamento Inteligência Artificial) e altera determinados atos legislativos da União.** 21 abr. 2021. Disponível em: https://eur-lex.europa.eu/resource.html?uri=cellar:e0649735-a372-11eb-9585-01aa75ed71a1.0004.02/DOC_1&format=PDF. Acesso em: 08 set. 2021.

¹⁶⁴ VASCONCELLOS, Carlos Eduardo. "A legislação europeia sobre IA poderá ter influência no mundo todo". **In: Consumidor Moderno**, 2021. Disponível em:

Como visto, já existe uma forte movimentação em contexto mundial para uma regulamentação específica em matéria de Inteligência Artificial, atualmente, a União Europeia vem se mostrando a grande pioneira na construção dessas regras, tendo como foco principal garantir que se tenha uma Inteligência Artificial confiável a partir dos parâmetros éticos, e que possa assegurar a prevalência dos Direitos Humanos e dos Direitos Fundamentais.

Além do destaque reconhecido no contexto mundial, a presente pesquisa pretende usar as orientações éticas para uma IA de Confiança como parâmetro para a regulamentação brasileira, pois, a cultura do Brasil tem mais proximidade com a cultura da União Europeia, por conta da base no bem-estar social¹⁶⁵.

2.3.2 Passos em direção da regulamentação sobre o uso da Inteligência Artificial no Brasil

Existem poucos projetos de lei em trâmite no Brasil que têm como foco principal a Inteligência Artificial. Atualmente, pode-se verificar a existência de cinco projetos que apresentam a IA como temática principal, e diversos outros que citam a tecnologia em seus artigos.

Como propostas do Senado Federal existe o Projeto de Lei 5.051/2019¹⁶⁶, de 16 de setembro de 2019, do senador Styvenson Valentim (PODEMOS/RN), que estabelece os princípios para o uso da Inteligência Artificial no Brasil. O Projeto de Lei 5.691/2019¹⁶⁷, de 25 de outubro de 2019, também do senador Styvenson, que institui a

<https://www.consumidormoderno.com.br/2021/07/02/legislacao-europeia-ia-influencia-mundo/>. Acesso em: 08 set. 2021.

¹⁶⁵ Sobre a compreensão de bem-estar social, Bento disserta que a função do Estado é de proteção social nos âmbitos educacional, trabalhista e do pleno emprego, da assistência à saúde e à seguridade, da prosperidade econômica e outros. Porém, vale destacar ainda que a providência —prestada a todo cidadão no sentido de lhe garantir uma renda mínima se dá não a título de caridade pública, mas de um direito, considerado fundamental para o bem-estar da vida em sociedade. Portanto a maior ingerência pública se justifica na indispensabilidade de uma autoridade legal que garanta a segurança e o amparo frente à vulnerabilidade que o país apresenta aos riscos de instabilidade nacional e internacional, tendo em vista o bem comum. Dessa forma, ao Estado cabe o planejamento e a implementação de políticas que assegurem a execução dos direitos supracitados visando a equidade de oportunidades. (BENTO, Leondardo V. **Governança e Governabilidade na reforma do Estado: entre eficiência e democratização**. Barueri: Manole, 2003, p. 7)

¹⁶⁶ BRASIL. Senado Federal. **Projeto de Lei n. 5051, de 2019**. Estabelece os princípios para o uso da Inteligência Artificial no Brasil. Disponível em: <https://legis.senado.leg.br/sdleg-getter/documento?dm=8009064&ts=1630421610171&disposition=inline>. Acesso em: 08 set. 2021.

¹⁶⁷ BRASIL. Senado Federal. **Projeto de Lei n. 5691, de 2019**. Institui a Política Nacional de Inteligência Artificial. Disponível em: <https://legis.senado.leg.br/sdleg-getter/documento?dm=8031122&ts=1630421785830&disposition=inline>. Acesso em: 08 set. 2021.

Política Nacional de Inteligência Artificial. E, o Projeto de Lei 872/2021¹⁶⁸, de 12 março de 2021, do senador Veneziano Rêgo (MDB – PB), que dispõe sobre o uso da Inteligência Artificial no Brasil.

Recentemente aprovado pela Câmara dos Deputados, o Projeto de Lei n. 21/2020¹⁶⁹, de 04 de fevereiro de 2020, do deputado Eduardo Bismarck (PDT – CE), que tem o objetivo de estabelecer princípios, direitos e deveres para o uso de Inteligência Artificial no Brasil, apresentando um conteúdo mais completo que as propostas do Senado, contudo, a legislação não é suficientemente profunda e se espera que o Senado Federal, no qual o projeto encontra-se em tramitação, melhore essas questões. Bem como existe o Projeto de Lei 240/2020¹⁷⁰, de 11 de fevereiro de 2020, do deputado Léo Moraes (PODE/RO), que cria a Lei da Inteligência Artificial, já apensado ao Projeto de Lei n. 21/2020, mas com um conteúdo mais simplório sobre o tema. Há um detalhe neste último que merece atenção: é o de justificativa mais bem fundamentada, relatando um panorama mundial da regulamentação da IA e novas tecnologias.

Contudo, duas normatizações, criadas pelo Conselho Nacional de Justiça merecem atenção, a Resolução n. 332/2020¹⁷¹, que “dispõe sobre a ética, a transparência e a governança na produção e no uso de Inteligência Artificial no Poder Judiciário”, e a Portaria n. 271/2020, que “regulamenta o uso de Inteligência Artificial no âmbito do Poder Judiciário”. Ambas apresentam um trabalho mais minucioso do que vem sendo desenvolvido pelo legislativo federal.

De acordo com a Resolução n. 332/2020, a IA aplicada aos sistemas judiciários visa a melhorar a eficiência e a qualidade da justiça, devendo respeitar os direitos fundamentais do indivíduo previstos na Convenção de Direitos Humanos e Convenção de Proteção de Dados Pessoais, levando em consideração também princípios como

¹⁶⁸ BRASIL. Senado Federal. **Projeto de Lei n. 872, de 2021**. Dispõe sobre o uso da Inteligência Artificial. Disponível em: <https://legis.senado.leg.br/sdleg-getter/documento?dm=8940096&ts=1634324707816&disposition=inline>. Acesso em: 08 set. 2021.

¹⁶⁹ BRASIL. Câmara dos Deputados. **Projeto de Lei n. 21, de 2020**. Estabelece princípios, direitos e deveres para o uso de inteligência artificial no Brasil, e dá outras providências. Disponível em: https://www.camara.leg.br/proposicoesWeb/prop_mostrarintegra;jsessionid=node015en4zcdsiu5ggqc-hdxtw9qjv11397236.node0?codteor=1853928&filename=PL+21/2020. Acesso em: 08 set. 2021.

¹⁷⁰ BRASIL. Câmara dos Deputados. **Projeto de Lei n. 240, de 2020**. Cria a Lei da Inteligência Artificial, e dá outras providências. Disponível em: https://www.camara.leg.br/proposicoesWeb/prop_mostrarintegra?codteor=1857143&filename=PL+240/2020. Acesso em: 08 set. 2021.

¹⁷¹ BRASIL. Conselho Nacional de Justiça. **Resolução nº 332 de 21 de agosto de 2020**. Dispõe sobre a ética, a transparência e a governança na produção e no uso de Inteligência Artificial no Poder Judiciário e dá outras providências. Disponível em: <https://atos.cnj.jus.br/atos/detalhar/3429>. Acesso em: 08 set. 2021.

a não discriminação, qualidade e segurança, transparência, imparcialidade e justiça, e o princípio “sob controle do usuário”.¹⁷²

Segundo a resolução, a implementação da IA contribui para a coerência dos processos e a agilidade de tramitação até a tomada de decisão, desde que observada a compatibilidade com os Direitos Fundamentais e critérios éticos, ressaltando que:

As decisões judiciais apoiadas pela Inteligência Artificial devem preservar a igualdade, a não discriminação, a pluralidade, a solidariedade e o julgamento justo, com a viabilização de meios destinados a eliminar ou minimizar a opressão, a marginalização do ser humano e os erros de julgamento decorrentes de preconceitos.¹⁷³

Contudo, há de ser feita uma observação relevante. É sabido que a resolução é conceituada como um ato legislativo de conteúdo concreto com efeitos internos, ou seja, em uma escala hierárquica comparativa entre resolução e lei, esta última é de relevância maior que a primeira.

Desta forma, Marcus Lívio explicou:

Posso assegurar que foi um documento amplamente debatido internamente pelo grupo de juízes. Mais de 30 colegas colaboraram para que esse ato normativo fosse consolidado, tivemos também a ajuda de muitos outros atores externos, **mas é importante deixar claro que não é um ato definitivo**. Estaremos sempre aptos a propor alternativas de solução ou melhor redação, sempre que instados, principalmente, pelo centro de estudos liderado pelo Ministro Luis Felipe Salmão.¹⁷⁴

Sendo assim, apesar de alguns avanços na área jurídica sobre o uso da inteligência artificial que veio com a publicação da Resolução n. 332, todavia, ainda existe a necessidade de uma regulamentação com força maior, ou seja, existe a necessidade de uma Lei própria, capaz de transmitir confiança e segurança a sociedade a partir do respeito a seus direitos fundamentais.

Tendo em vista a necessidade de se estabelecer uma regulamentação específica sobre a inteligência artificial, mostra-se relevante mencionar que vem ganhando um certo destaque por recentemente já ter passado pela mesa de votação e ter sido aprovado com alterações no Plenário e encaminhado para votação no Senado Federal, em 29

¹⁷² BENTES, Dorinethe dos Santos; CURVINO, Sarah Fernandes. O acesso à Justiça e a Inteligência Artificial: uma análise da Resolução n.º 332/2020 do Conselho Nacional de Justiça. **In: XI Congresso RECAJ-UFGM**, Belo Horizonte, 2020.

¹⁷³ BRASIL. Conselho Nacional de Justiça. **Resolução nº 332 de 21 de agosto de 2020**. Dispõe sobre a ética, a transparência e a governança na produção e no uso de Inteligência Artificial no Poder Judiciário e dá outras providências. Disponível em: <https://atos.cnj.jus.br/atos/detalhar/3429>. Acesso em: 08 set. 2021, p. 01.

¹⁷⁴ LÍVIO, Marcus. Um retrato sobre o uso da inteligência artificial no Poder Judiciário. **In: Justiça & Cidadania**. 2021. Disponível em: <https://editorajc.com.br/um-retrato-sobre-o-uso-da-inteligencia-artificial-no-poder-judiciario/>. Acesso em: 08 set. 2021.

de setembro de 2021, o Projeto de Lei nº 21/2020, de autoria do Deputado Eduardo Bismarck (PDT/CE), que visa estabelecer os princípios, direitos e deveres e instrumentos de governança para o uso da Inteligência Artificial, no qual traz como objetivo em sua ementa que: “estabelece princípios, direitos e deveres para o uso de IA no Brasil”. Salienta-se, no entanto, que o referido projeto será melhor analisado no terceiro capítulo, tendo em vista a necessidade de fazer um paralelo com a regulamentação internacional.

O texto de apresentação do PL 21/2020 além de apregoar o respeito aos direitos humanos e aos valores democráticos (art. 4, III) aponta para a transparência e explicabilidade dos sistemas de inteligências artificial (art. 6, IV) e a governança multiparticipativa entre os diferentes agentes envolvidos, públicos ou privados (art. 10, VI).

Importante salientar também como o legislador pretende abordar de forma expressa o desenvolvimento tecnológico, a inovação, o estabelecimento dos parâmetros da livre iniciativa e concorrência, o respeito aos direitos humanos e aos valores democráticos, bem como a privacidade e proteção de dados. Delineiam ainda, sobre o desenvolvimento da IA, da ética, livre de preconceitos, ponto de maior relevância quando se trata da implementação da IA. Estabelece o uso responsável de inteligência artificial no Brasil, bem como a sua finalidade, conforme prevê o art. 6º:

Artigo 6º. [...] I - finalidade: uso de inteligência artificial para buscar resultados benéficos para as pessoas do planeta, com o fim de aumentar as capacidades humanas, reduzir as desigualdades sociais e promover o desenvolvimento sustentável.¹⁷⁵

Em síntese, o Projeto de Lei n. 21/2020 tem quatro bases: primeiro, conceituação dos principais termos para regulamentação e criação de categorias de agentes com atribuições e responsabilidades específicas; segundo, instituição dos fundamentos, princípios e objetivos orientadores do uso responsável da Inteligência Artificial; terceiro, direitos e deveres de todos os envolvidos; e quarto, diretrizes para atuação do poder público. O projeto também prevê a criação de dois tipos de agentes de Inteligência Artificial: de desenvolvimento e o de operação, e lhes atribui deveres e responsabilidades específicas, observadas as suas respectivas funções.¹⁷⁶

¹⁷⁵ BRASIL. Câmara dos Deputados. **Projeto de Lei n. 21, de 2020**. Estabelece princípios, direitos e deveres para o uso de inteligência artificial no Brasil, e dá outras providências. Disponível em: https://www.camara.leg.br/proposicoesWeb/prop_mostrarintegra;jsessionid=node015en4zcdsiu5ggqc-hdxtw9qjv11397236.node0?codteor=1853928&filename=PL+21/2020. Acesso em: 08 set. 2021.

¹⁷⁶ TOSTES, Marcelo. Marco Legal da Inteligência Artificial traz benefícios para o Brasil. **In: Exame**, 2021. Disponível em: <https://exame.com/bussola/marco-legal-da-inteligencia-artificial-traz-beneficios->

Apesar de alguns passos já terem sido direcionados para uma regulamentação do uso da Inteligência Artificial no Brasil, enquanto não há legislação aprovada, a melhor recomendação é observar práticas internacionais, em especial aquelas da Organização para a Cooperação e Desenvolvimento Econômico (OCDE). Em relação à proteção de dados, a legislação proposta se soma à LGPD para mitigar riscos sobre os processados por tecnologias cognitivas.

3 UM ESTUDO SOBRE A CONSTRUÇÃO DA INTELIGÊNCIA ARTIFICIAL DE CONFIANÇA SOB O ENFOQUE DOS DIREITOS HUMANOS

Para que se possa considerar uma IA de confiança é necessário considerar os parâmetros definidos pelos Direitos Humanos, por ser a norma internacional que traz os parâmetros de dignidade humana, a qual retrata a conquista pela igualdade no que tange aos bens imateriais e materiais nas distintas maneiras de vida e de viver que existem.

3.1 O PROCESSO DE INTERNACIONALIZAÇÃO E UNIVERSALIZAÇÃO DOS DIREITOS HUMANOS

No que se refere ao conceito de direitos humanos, Dornelle¹⁷⁷ afirma que não há uma uniformidade conceitual clara, embora algumas pessoas teimem em apresentar uma única maneira de definir os direitos humanos, o seu conceito varia de acordo com a concepção político-ideológica que se tenha.

O que importa é que os direitos ou valores (dependendo da ótica) considerados fundamentais sofrem uma variação de acordo com o modo de organização da vida social. E, portanto, impossível a existência de uma única fundamentação dos direitos humanos.¹⁷⁸

No mesmo sentido, ainda que os direitos humanos não se limitam às declarações da Organização das Nações Unidas (ONU) e dos outros organismos internacionais, são de extrema importância para concretização no âmbito das disposições constitucionais dos Estados, de forma a constituir um conjunto de princípios norteados do direito internacional.¹⁷⁹

Os direitos humanos, se convertem em tema de legítimo interesse internacional, transcendendo o âmbito da jurisdição doméstica específica de cada Estado, embora não pretenda substituir o sistema nacional, o sistema internacional de proteção de direitos humanos possui natureza subsidiária, quando falham as

¹⁷⁷ DORNELLE, João Ricardo. **O que são direitos humanos**. São Paulo. 5. ed., São Paulo: Editora Brasiliense Ltda, 2013.

¹⁷⁸ DORNELLE, João Ricardo. **O que são direitos humanos**. São Paulo. 5. ed., São Paulo: Editora Brasiliense Ltda, 2013, p. 15-16.

¹⁷⁹ TOSI, Giuseppe. Os Direitos Humanos: Reflexões iniciais. **In:** TOSI, Giuseppe (coord.). **Direitos Humanos: história, teoria e prática**. João Pessoa: Editora UFPB, 2004. Disponível em: https://www.academia.edu/40425908/Direitos_Humanos_Teoria_e_Pratica_Giusepp_e_Tosi. Acesso em: 09 out. 2021.

instituições nacionais.¹⁸⁰

Sendo assim, tem por objetivo discorrer sobre o processo de internacionalização dos direitos humanos, a partir da universalidade trazida pela da Declaração Universal de Direitos Humanos (DUDH) de 1948, para depois, frente ao processo de desconstrução da teórica clássica dominante, apresentar breves considerações sobre a teoria crítica desenvolvida por Professor Herrera Flores.

3.1.1 O caráter universal da Declaração Universal de Direitos Humanos de 1948

Em meados do século XX, ocorreu o surgimento de três instituições responsáveis pela redefinição do conceito tradicional de soberania estatal absoluta, acabando por interferir na liberdade e autonomia dos Estados perante a comunidade global. “O Direito Humanitário, a Liga das Nações e a Organização Internacional do Trabalho situam-se como os primeiros marcos do processo de internacionalização dos direitos humanos”.¹⁸¹

O objetivo era de romper com a ideia tradicional que colocava os Estados como únicos sujeitos de direito internacional, de forma a incorporar uma visão mais humanitária, a partir da necessidade de elevar o indivíduo enquanto sujeito de direito internacional e não mais como um objeto subordinado a jurisdição doméstica de cada país.

[...] para que os direitos humanos se internacionalizassem, foi necessário redefinir o âmbito e o alcance do tradicional conceito de soberania estatal, a fim de permitir o advento dos direitos humanos como questão de legítimo interesse internacional. Foi ainda necessário redefinir o status do indivíduo no cenário internacional, para que tornasse verdadeiro sujeito de Direito Internacional.¹⁸²

No entanto, a internacionalização dos direitos humanos se consolidou, de fato, com a promulgação da Declaração Universal de Direitos Humanos de 1948, aprovada pela Resolução n. 217 A (III), da Assembleia Geral das Nações Unidas (AGNU), depois de o mundo ter experienciado as barbáries cometidas pela política nazista

¹⁸⁰ TOSI, Giuseppe. Os Direitos Humanos: Reflexões iniciais. **In:** TOSI, Giuseppe (coord.). **Direitos Humanos: história, teoria e prática.** João Pessoa: Editora UFPB, 2004. Disponível em: https://www.academia.edu/40425908/Direitos_Humanos_Teoria_e_Pratica_Giusepp_e_Tosi. Acesso em: 09 out. 2021.

¹⁸¹ PIOVESAN, Flávia. **Direitos humanos e o direito constitucional internacional.** 14. ed., São Paulo: Saraiva, 2013.

¹⁸² PIOVESAN, Flávia. **Direitos humanos e o direito constitucional internacional.** 14. ed., São Paulo: Saraiva, 2013, p. 183.

alemã, no conflito militar global, conhecido por Segunda Guerra Mundial.

O fundamento base da DUDH, surgiu pela necessidade de reconstrução dos direitos humanos sob a ótica de valorização da pessoa humana, considerando a destruição física, psíquica e social do indivíduo, causada pelo totalitarismo da era Hitler.

As violações a nível mundial eram tamanhas que “apresentando o Estado como o grande violador de direitos humanos, a Era Hitler for marcada pela lógica da destruição e da descartabilidade da pessoa humana o que resultou no extermínio de onze milhões de pessoas”.¹⁸³

Em meio aos conflitos políticos, econômicos, sociais e culturais da época, bem como a inefetividade dos sistemas internacionais de proteção dos direitos humanos vigentes até então, surgiu a criação de uma declaração capaz de evitar uma possível terceira guerra mundial e promover a paz duradoura entre todos os povos, de modo a reconstruir os direitos humanos.

Nesse contexto, desenha-se o esforço de reconstrução dos direitos humanos, como paradigma e referencial ético a orientar a ordem internacional contemporânea. Se a Segunda Guerra significou a ruptura com os direitos humanos, o pós-guerra deveria significar a sua reconstrução. [...] A necessidade de uma ação internacional mais eficaz para a proteção dos direitos humanos impulsionou o processo de internacionalização desses direitos, culminando na criação da sistemática normativa de proteção internacional, que faz possível a responsabilização do Estado no domínio internacional quando as instituições nacionais se mostram falhas ou omissas na tarefa de proteger os direitos humanos.¹⁸⁴

Segundo Tosi¹⁸⁵ a Declaração Universal serviu para reafirmar a concepção clássica jus naturalista, de forma a promover e homenagear a tradição dos direitos naturais. Ainda, se nos atentarmos a própria redação do artigo primeiro da DUDH, veremos a intenção dos redatores em repetir as três palavras de ordem da Revolução Francesa de 1789 (liberdade, igualdade e fraternidade).

A Declaração afirma que “todas as pessoas nascem livres e iguais”; esta formulação é uma citação explícita da “Declaração dos direitos do homem e do cidadão” da Revolução Francesa. Ela quer significar o caráter natural dos direitos, enquanto inerentes à natureza de cada ser humano, pelo reconhecimento de sua intrínseca dignidade. Neste sentido, os direitos tornam-se um conjunto de valores éticos universais que estão “acima” do

¹⁸³ PIOVESAN, Flávia. **Direitos humanos e o direito constitucional internacional**. 14. ed., São Paulo: Saraiva, 2013, p. 190.

¹⁸⁴ PIOVESAN, Flávia. **Direitos humanos e o direito constitucional internacional**. 14. ed., São Paulo: Saraiva, 2013, p. 190-191.

¹⁸⁵ TOSI, Giuseppe. Os Direitos Humanos: Reflexões iniciais. *In*: TOSI, Giuseppe (coord.). **Direitos Humanos: história, teoria e prática**. João Pessoa: Editora UFPB, 2004. Disponível em: https://www.academia.edu/40425908/Direitos_Humanos_Teoria_e_Pratica_Giusepp_e_Tosi. Acesso em: 09 out. 2021.

nível estritamente jurídico e que devem orientar a legislação dos Estados.¹⁸⁶

Ainda de acordo com este autor “isto foi fruto de uma negociação entre os dois grandes blocos do pós-guerra, o bloco socialista – que defendia os direitos econômicos e sociais – e o bloco capitalista – que defendia os direitos civis e políticos”.¹⁸⁷

A DUDH foi marcada pela tentativa de uniformização de entendimentos entre as duas principais vertentes políticas contemporâneas ocidentais da época, ou seja, de um lado, os direitos civis e políticos, típicos da tradição contratualista liberal e do outro, os direitos econômicos, sociais e culturais, intrínsecos a tradição socialista.¹⁸⁸

Piovesan explica que “à luz da perspectiva histórica, observa-se que até então intensa era a dicotomia entre o direito à liberdade e o direito à igualdade”.¹⁸⁹

Dessa forma, o documento objetivou encontrar consensos, mediante valores universais, que pudessem servir de orientação para todos os países membros e sendo aplicável a todas as pessoas, independentemente de cor, raça, sexo e religião, para a Declaração Universal a condição de pessoa é o requisito único e exclusivo para a titularidade de direitos.¹⁹⁰

Sobretudo, pela ideia de indivíduo enquanto membro direto da comunidade mundial, cidadão do mundo, protegido internacionalmente, pelo fundamento maior de respeito à dignidade humana.

A reação à barbárie do nazismo e dos fascismos em geral levou, no pós-guerra, à consagração da dignidade da pessoa humana no plano internacional e interno como valor máximo dos ordenamentos jurídicos e princípio orientador da atuação estatal e dos organismos internacionais.¹⁹¹

¹⁸⁶ TOSI, Giuseppe. Os Direitos Humanos: Reflexões iniciais. *In*: TOSI, Giuseppe (coord.). **Direitos Humanos: história, teoria e prática**. João Pessoa: Editora UFPB, 2004. Disponível em: https://www.academia.edu/40425908/Direitos_Humanos_Teoria_e_Pratica_Giusepp_e_Tosi. Acesso em: 09 out. 2021, p. 19.

¹⁸⁷ TOSI, Giuseppe. Os Direitos Humanos: Reflexões iniciais. *In*: TOSI, Giuseppe (coord.). **Direitos Humanos: história, teoria e prática**. João Pessoa: Editora UFPB, 2004. Disponível em: https://www.academia.edu/40425908/Direitos_Humanos_Teoria_e_Pratica_Giusepp_e_Tosi. Acesso em: 09 out. 2021, p. 16.

¹⁸⁸ TOSI, Giuseppe. Os Direitos Humanos: Reflexões iniciais. *In*: TOSI, Giuseppe (coord.). **Direitos Humanos: história, teoria e prática**. João Pessoa: Editora UFPB, 2004. Disponível em: https://www.academia.edu/40425908/Direitos_Humanos_Teoria_e_Pratica_Giusepp_e_Tosi. Acesso em: 09 out. 2021.

¹⁸⁹ PIOVESAN, Flávia. **Direitos humanos e o direito constitucional internacional**. 14. ed., São Paulo: Saraiva, 2013, p. 211.

¹⁹⁰ PIOVESAN, Flávia. **Direitos humanos e o direito constitucional internacional**. 14. ed., São Paulo: Saraiva, 2013, p. 210.

¹⁹¹ BARCELLOS, Ana Paula de. Normatividade dos princípios e o princípio da dignidade da pessoa humana na Constituição de 1988. *In*: **Revista de Direito Administrativo**, Rio de Janeiro, vol. 221, p. 159-188, jul./set. 2000. Disponível em: <http://bibliotecadigital.fgv.br/ojs/index.php/rda/article/view/47588/45167>. Acesso em: 09 out. 2021, p. 162.

Neste sentido, a DUDH buscou defender os direitos dos ostensivamente mais fracos, num mundo frágil, hostil e repleto de desigualdades, consagrou em seu primeiro artigo o seguinte:

Artigo 1º Todos os seres humanos nascem livres e iguais em dignidade e direitos. São dotados de razão e consciência e devem agir em relação uns aos outros com espírito de fraternidade.¹⁹²

A máxima de que todos os seres humanos nascem livres e iguais em dignidade e direitos, além de categorizar o indivíduo como pertencente a toda comunidade global, nos remete a ideia de que a igualdade é algo intrínseco a própria natureza de ser de cada indivíduo, pelo simples fato de existir.

No âmbito do direito internacional, a trajetória da dignidade humana inicia com a sua expressa previsão na Constituição da OIT (1946), mas se afirma – como valor jurídico universal – mediante a Declaração dos Direitos Humanos da ONU, de 1948 (inclusive no Preâmbulo), sendo tematizada no âmbito do Preâmbulo (que refere a dignidade inerente a todos os membros da família humana) e no artigo 1º (todos os seres humanos nascem livres e iguais em dignidade) [...].¹⁹³

Portanto, até aqui, foi possível identificar os esforços da comunidade internacional para a consolidação jurídica de uma base mínima de direitos, capaz de alcançar a todos os indivíduos do planeta terra, a partir da ideia abstrata de que “ser humano” é sinônimo de “ter direitos”.

3.1.2 Teoria crítica dos Direitos Humanos sob o olhar de Herrera Flores

Na segunda década do século XX (mais especificamente entre o período do fim da Primeira Guerra Mundial ao fim da Segunda Guerra Mundial), filósofos, artistas, sociólogos e pensadores se uniram, inicialmente, a fim de discutir e pesquisar sobre a teoria marxista e sua aplicação na sociedade da época. Max Horkheimer (1895-1973) e Theodor W. Adorno (1903- 1969) – juntamente com intelectuais teóricos que vieram a passar por este Instituto durante aquele século – tomaram a linha de frente deste projeto, no qual se localizava na Escola de Frankfurt, Instituto de Pesquisas Sociais (Institut für Sozialforschung) da Universidade de Frankfurt, estudo que em

¹⁹² ORGANIZAÇÃO DAS NAÇÕES UNIDAS. **Declaração Universal dos Direitos Humanos**. 1948. Disponível em: <https://www.oas.org/dil/port/1948%20Declara%C3%A7%C3%A3o%20Universal%20dos%20Direitos%20Humanos.pdf>. Acesso em: 09 out. 2021.

¹⁹³ SARLET, Ingo Wolfgang. **Dignidade (da pessoa) humana e direitos fundamentais na Constituição Federal de 1988**. 10. ed., Porto Alegre: Livraria do Advogado Editora, 2019, p. 76.

breve seria nomeado o estudo da “teoria crítica”.¹⁹⁴

Mogendorff¹⁹⁵, explana que os estudos da teoria social crítica tiveram como pilar, além do marxismo, os ideais defendidos por Sigmund Freud (1856- 1939) e Friedrich Nietzsche (1844-1900), buscando compreender as tensões sociais causadas pelo sistema capitalista e suas respectivas motivações para um modo de vida aceitável e desejado.

Segundo Gallardo¹⁹⁶, “uma teoria é crítica quando busca compreender tanto as situações de discriminação, como a base sistêmica ou estrutural que as produz”¹⁹⁷.

O pensamento crítico não se baseia apenas em uma situação onde a crítica está presente. Mais do que isso, ele engloba em sua linha de raciocínio as tensões sociais e as maneiras de se chegar ao objetivo principal: a emancipação humana, de maneira a distanciar a forma de pensamento decolonial da sociedade. Carballido explica:

Dito isso, é importante ressaltar que nem todo o pensamento que critica algo, pelo simples fato de criticar, é pensamento crítico. A particularidade da crítica no pensamento crítico, e esta deve ser evidente em seus conteúdos, reside em um determinado ponto de vista a partir do qual a crítica é realizada no ponto de vista da emancipação humana.¹⁹⁸

Ainda assim, a teoria crítica abarcada pelo pensamento crítico dispõe de algumas características significativas para tornar-se uma teoria afirmativa e que de fato entenda e discuta as discrepâncias sociais.

O viés crítico explorado por teóricos da Escola de Frankfurt no que tange à uma visão diferenciada do mundo moderno inspirou o jurista, professor e defensor dos direitos humanos Joaquín Herrera Flores ao desenvolver uma longa pesquisa dentro do campo de estudo destes direitos, onde ele argumenta de diversas formas as falhas no discurso e na reprodução humanista com o objetivo – não imediato – de alcançar

¹⁹⁴ MOGENDORFF, Janine Regina. A Escola de Frankfurt e seu legado. *In: Verso e Reverso*, Porto Alegre, vol. 26, n. 63, p. 152-159, set. 2012, p. 152-153.

¹⁹⁵ MOGENDORFF, Janine Regina. A Escola de Frankfurt e seu legado. *In: Verso e Reverso*, Porto Alegre, vol. 26, n. 63, p. 152-159, set. 2012, p. 154.

¹⁹⁶ GALLARDO, Helio. Teoría crítica y derechos humanos: una lectura latinoamericana. *In: Revista de Derechos Humanos y Estudios Sociales*, São Luís Potosí, vol. 2, n. 4, p. 57-89, jul. 2010. Disponível em: <http://www.derecho.uaslp.mx/Documents/Revista%20REDHES/N%C3%BAmero%204/Redhes4-03.pdf>. Acesso em: 30 set. 2021, p. 68, tradução nossa.

¹⁹⁷ Texto original: “Una teoría es crítica si busca comprender tanto las situaciones de discriminación como la base sistémica o estructural que las produce”.

¹⁹⁸ Texto original: “Dicho esto, vale la pena señalar que no todo pensamiento que critica algo, por el mero hecho de criticar, es pensamiento crítico. La particularidad de la crítica en el pensamiento crítico, y ello ha de ser evidente en sus contenidos, reside en un determinado punto de vista a partir del cual dicha crítica es realizada, siendo este punto de vista el de la emancipación humana”. (CARBALLIDO, Manuel E. Gándara. *Los derechos humanos en el siglo XXI: una mirada desde el pensamiento crítico*. Buenos Aires: CLACSO, 2019, p. 25, tradução nossa)

a igualdade perante a dignidade humana, bem como aponta meios de legitimar tal discurso de uma maneira que haja resultados efetivos à sociedade.

Sua crítica abrange o modelo de sistema econômico, social e cultural usado pelo mundo ocidental, a utilização da teoria tradicional em áreas que pedem o pensamento crítico e o desenvolver das relações humanas não como formas de raciocínios prontos, mas de desenvolvimento de teorias e práticas distintas ao que já é feito, e o cenário em que os direitos humanos foram declarados e efetivados.¹⁹⁹

Em sua obra intitulada “Teoria crítica dos direitos humanos: os direitos humanos como produtos culturais”, Flores expõe sobre uma teoria crítica afirmativa e forte para ser utilizada no âmbito dos direitos humanos. De início, para que seja efetivo tal discurso, ele cita três tendências a serem superadas “para se chegar a uma teoria crítica que realmente nos sirva para nos diferenciar das concepções hegemônicas dos direitos humanos”, sendo a primeira ideia de que não há como ter um pensamento crítico que abranja e generalize todas as pessoas e universalize os próprios direitos em questão, visto que cada sociedade em si possui suas particularidades, as quais não devem ser vistas de um modo geral como se todos fossem semelhantes. Isso não significa que todas as formas gerais de compreensão e proteção dos direitos humanos conquistadas até hoje devem ser desconsideradas (segunda ideia). Pelo contrário: a criação das Nações Unidas, sua respectiva declaração e seus conselhos, organizações, comissões e tribunais, tal qual o nascimento de ONGs e sindicatos de diferentes lugares são, de fato, essenciais para o desenvolvimento e a propagação da ideia central destes direitos, fundamentados na premissa de garantia da dignidade humana. Por último, Flores entende que a força de uma teoria crítica dos direitos humanos não deve ser medida através de seu efeito na sociedade, sendo ele no presente ou no futuro. Isto porque o teórico, considerando-o em seu meio acadêmico, não possui as práticas e vivências suficientes para ditar o que e como elas devem se comportar, ressaltando que a teoria deve sempre se adaptar à realidade, nunca o contrário. Ademais, a ideia de uma teoria crítica dos direitos humanos não é traçar caminhos futuros e irreais de comportamento humano, mas entender as tensões sociais do presente momento e construir um caminho alternativo àquele que nos é conhecido através do regime capitalista e do conhecimento geral dos direitos

¹⁹⁹ FLORES, Joaquín Herrera. **A (re)invenção dos direitos humanos**. Florianópolis: Fundação Boiteux, 2009.

humanos, a fim de buscar um resultado mais próxima da emancipação humana.²⁰⁰ De acordo com os ideais de Boaventura de Sousa Santos, o professor em questão comenta sobre:

[...] há que se ampliar e expandir as lutas que se dão no presente e reduzir o plano das expectativas futuras. Único caminho que nos possibilitará ampliar e expandir nossas experiências cotidianas e, a partir delas, ir construindo e reconstruindo alternativas reais e possíveis no mundo em que vivemos.²⁰¹

O conceito de universalização de ideias, pensamentos, do “real” e do “racional” por volta do século XVI, é criticado por Flores quando ele comunica que estas questões não poderiam ser universais e ditas como uma visão de mundo única, uma vez que, na realidade, é a visão de um grupo pequeno de pessoas que, de acordo com a história, vem de uma classe opressora e colonizadora, onde sua grande maioria é composta por homens brancos os quais decidiram que seu racional seria o certo, excluindo maneiras diferentes de ver o mundo em sua essência plural e cultural.²⁰²

Esta forma de universalizar uma realidade e uma noção de mundo como algo certo é oriunda do que hoje conhecemos como o mundo Ocidental. O autor em discussão entende que o Ocidente pode ser, assim, considerado também por questões econômicas, geoestratégicas, reacionárias ou revolucionárias, no entanto a maneira de ter seu próprio pensamento como o único e o correto, excluindo os demais também contribui para a conceituação do termo.²⁰³

Isso implica na forma como os direitos humanos são vistos e entendidos, já que o seu sistema de criação veio de mentes pensantes ocidentais. Tem-se como exemplo: apesar de a Revolução Francesa de fato revolucionar a população de seu país e de outros com seus ideais “liberdade, igualdade e fraternidade” por meio da concepção de universalização dos direitos resguardados na declaração escrita e oriundos da revolução, ao se enxergar da maneira proposta (universal) por eles, vê-se que não foi isso que aconteceu. O Haiti, país caribenho colonizado pela França, não recebeu este “tratamento” defendido pela revolução da nação europeia em questão, muito menos conseguiu garantir os direitos ali discutidos aos seus cidadãos, caindo por terra a ideia “universal” dos princípios do acontecimento francês. Flores

²⁰⁰ FLORES, Joaquín Herrera. **Teoria crítica dos direitos humanos**: os direitos humanos como produtos culturais. Rio de Janeiro: Editora Lumen Juris, 2009, p. 36-38.

²⁰¹ FLORES, Joaquín Herrera. **Teoria crítica dos direitos humanos**: os direitos humanos como produtos culturais. Rio de Janeiro: Editora Lumen Juris, 2009, p. 38.

²⁰² FLORES, Joaquín Herrera. **Teoria crítica dos direitos humanos**: os direitos humanos como produtos culturais. Rio de Janeiro: Editora Lumen Juris, 2009, p. 38-42.

²⁰³ FLORES, Joaquín Herrera. **Teoria crítica dos direitos humanos**: os direitos humanos como produtos culturais. Rio de Janeiro: Editora Lumen Juris, 2009, p. 38.

cita o assunto no terceiro paradoxo dos direitos humanos universais, conceituado “paradoxo de duplo critério”:

Situemo-nos na França da burguesia revolucionária de 1789 e deleitemo-nos com a famosa trindade dos valores “liberdade, igualdade e fraternidade”, proclamados como valores universais pela própria razão ilustrada. Vamos dar um salto oceânico e visitar o Haiti de finais do século XVIII. Tão-somente uns meses depois da famosa “declaração de direitos”, os escravos e recém libertos haitianos rebelaram-se contra a ordem escravista que os tinha submetido à mais execrável das dominações. Os escravos e libertos levantaram suas armas contra os escravistas e os proprietários de fazendas, exigindo os mesmos valores racionais que a burguesia “revolucionária” francesa logrou converter em norma: liberdade, igualdade e fraternidade.²⁰⁴

Outrossim, ainda sobre o mesmo contexto histórico, não era preciso ir a outro continente para entender que é contraditado tal conceito de universalização. Na própria França os direitos conquistados não se aplicavam a todos, como por exemplo, às mulheres, sendo este mais um motivo pelo qual não se deve cair na tentação de acreditar que os direitos humanos são de fato universais como havia sido dito e como até hoje é entendido. Flores também comenta, agora sobre a luta e persistência da feminista francesa Olimpe de Gouges ao tentar abranger as mulheres aos direitos ora declarados:

A execução pública de Olimpe de Gouges, uma ousada e temerária mulher que se apresenta diante da Assembleia do Povo francês exigindo que a Revolução também levasse em conta as mulheres, nos mostra que aquela “declaração racional” tampouco era para todas. A guilhotina encarregou-se de calar as reivindicações antipatriarcais que Olimpe propôs que se incluíssem no “racional”, formulando por escrito – oh, afronta! – sua tese em uma contradecaração de direitos das mulheres e das cidadãs. A lâmina afiada que cortou a cabeça da girondina Olimpe de Gouges calou durante muitas décadas as reivindicações das mulheres, as quais tiveram que aportar, como comentamos, centenas de vítimas à luta pela consecução do mero direito a votar. As mulheres, perguntava-se Virgínia Wolf, não são racionais? Por que existe esse duplo critério que mede de um modo diferente a racionalidade masculina e a irracionalidade feminina?²⁰⁵

A universalização dos direitos humanos não procede por conta do fator apresentado acima e em razão do fator cultural, onde não se pode simplesmente impor os direitos humanos em povos e culturas que não os tem como parte de sua vivência, e que vivem sob diferentes modos de socialização e cultura. Flores, em sua obra intitulada “A (re)invenção dos direitos humanos”, menciona um exemplo do último fator:

É fácil ver a complexidade dos direitos, pois em grande quantidade de ocasiões tentam se impor em face de concepções culturais que nem sequer

²⁰⁴ FLORES, Joaquín Herrera. **Teoria crítica dos direitos humanos**: os direitos humanos como produtos culturais. Rio de Janeiro: Editora Lumen Juris, 2009, p. 55.

²⁰⁵ FLORES, Joaquín Herrera. **Teoria crítica dos direitos humanos**: os direitos humanos como produtos culturais. Rio de Janeiro: Editora Lumen Juris, 2009, p. 56.

têm sua bagagem linguística o conceito de direito (como é o caso e inumeráveis cosmovisões de povos e nações indígenas). Isso gera graves conflitos de interpretação em relação aos direitos humanos que se deve saber gerir sem imposições nem colonialismos.²⁰⁶

Na busca pelo entendimento de como os direitos humanos deveriam ser tratados e expostos, Flores²⁰⁷, por meio do pensamento crítico, aposta em duas complexidades que se confundem: a complexidade empírica e a complexidade jurídica. Na primeira, o professor destaca o conceito de que o empírico compete ao “ter direitos”, enquanto o normativo – ou jurídico; consiste na ideia de “o que/como deveríamos ter, através do direito”. O fato de o art. 5º, caput, da Constituição Federal de 1988²⁰⁸ dispor que “todos são iguais perante a lei [...]”, não significa que de fato todos os seres humanos são iguais. O uso da norma serve para definir o que deve ser e não o que realmente é, distinguindo-se, portanto, do conhecimento empírico e do pensamento de que só porque é norma ou está na lei, é.

Dessa maneira, quando a Declaração Universal dos Direitos Humanos (1948) cita e descreve todos os direitos das pessoas, não há como se interpretar de que o que lá está escrito já acontece/é, mas o que deveria acontecer/ser. Flores comenta sobre as complexidades empírica e jurídica e as exemplifica de maneira brilhante:

[...] Ao confundir o empírico com o normativo, parece que os direitos estão desde o primeiro momento conseguidos e incluídos na vida concreta das pessoas. Todavia, com apenas uma olhada ao nosso redor vemos que isso não ocorre assim. Utilizando unicamente os Relatórios Anuais do prestigiado Programa das Nações Unidas para o Desenvolvimento (PNUD), vemos, muitas vezes espantados, como o abismo entre os países pobres e os países ricos estão se criando bolsões de pobreza e desemprego, ante os quais as teorias econômicas e jurídicas não podem, ou não querem, reagir. **E, no entanto, se segue dizendo, talvez com boa vontade, que todos “têm” os mesmos direitos pelo simples fato de ter nascido. Ter nascido onde?**²⁰⁹

Continuamente citando-o, Flores²¹⁰ dedica uma parte do capítulo dois de sua obra trazendo quatro condições e cinco deveres básicos a fim de se ter uma teoria de fato crítica, portanto, pragmática dos direitos em estudo.

Quanto às quatro condições, a primeira versa sobre se ter uma visão realista

²⁰⁶ FLORES, Joaquín Herrera. **A (re)invenção dos direitos humanos**. Florianópolis: Fundação Boiteux, 2009, p. 43.

²⁰⁷ FLORES, Joaquín Herrera. **A (re)invenção dos direitos humanos**. Florianópolis: Fundação Boiteux, 2009, p. 43-48.

²⁰⁸ BRASIL. [Constituição, 1988]. **Constituição da República Federativa do Brasil de 1988**. Disponível em: http://www.planalto.gov.br/ccivil_03/constituicao/constituicao.htm. Acesso em: 17 out. 2021.

²⁰⁹ FLORES, Joaquín Herrera. **A (re)invenção dos direitos humanos**. Florianópolis: Fundação Boiteux, 2009, p. 48, grifo nosso.

²¹⁰ FLORES, Joaquín Herrera. **A (re)invenção dos direitos humanos**. Florianópolis: Fundação Boiteux, 2009, p. 60-68.

do mundo e diferente da qual é vista pelo senso comum. O pessimismo de certa forma atrapalha o exercício da teoria crítica, porquanto aquele é carregado de uma noção não real de situações da vida, tornando-as irreversíveis, enquanto que o pensamento crítico se baseia justamente no contrário, em “toda realidade como suscetível de quebra e transformação”. Ainda, a primeira condição aborda a importância do entendimento concreto e existente da sociedade, sobretudo de suas imprecisões e falhas, para chegar numa concepção de realidade onde todos estão e onde se quer chegar/estar para alcançar o objetivo principal: a dignidade humana.²¹¹

Na segunda condição, o autor em tela reforça que o pensamento crítico é um pensamento de combate, por isso, carece de compreensão ao identificar o adversário e firmar os objetivos e as intenções nas lutas, para que estas sejam inseridas na sociedade e haja mobilização. Também é citado neste bloco o empoderamento cidadão, cuja definição é justamente cenários nos quais movimentos sociais, distintos dos habitados e prevalecentes, ocupam um espaço maior nestes locais através de suas próprias linguagens e formas de falar sobre pautas específicas, concebendo-a Flores como algo “muito importante para uma teoria crítica dos direitos humanos”.²¹²

No que tange à terceira condição para ter uma teoria realista sobre os direitos, é mencionado de início a necessidade de reconhecer que o pensamento crítico é criado, em seu conceito majoritário, como pretexto obrigatório em e para grupos, hoje considerados “minorias” a fim de auxiliá-los na chegada do objetivo principal deste pensamento de forma segura e exitosa. Ademais, Flores²¹³ opina que as críticas social, artística e cultural devem se unir, e relembra que o direito é um produto cultural e normas jurídicas não nasceram das mãos de um legislador, mas de uma sociedade que as criou com o intuito de alcançar o acesso aos bens comuns e materiais defendidos, hoje, pelo sistema capitalista. Flores reforça:

Uma teoria crítica do direito deve se sustentar, então, sobre dois pilares: o reforço das garantias formais reconhecidas juridicamente, mas, igualmente, o empoderamento dos grupos mais desfavorecidos ao lutar por novas formas, mais igualitárias e generalizadoras, de acesso aos bens protegidos pelo direito.²¹⁴

²¹¹ FLORES, Joaquín Herrera. **A (re)invenção dos direitos humanos**. Florianópolis: Fundação Boiteux, 2009, p. 61.

²¹² FLORES, Joaquín Herrera. **A (re)invenção dos direitos humanos**. Florianópolis: Fundação Boiteux, 2009, p. 62.

²¹³ FLORES, Joaquín Herrera. **A (re)invenção dos direitos humanos**. Florianópolis: Fundação Boiteux, 2009.

²¹⁴ FLORES, Joaquín Herrera. **A (re)invenção dos direitos humanos**. Florianópolis: Fundação Boiteux, 2009, p. 65.

Na última condição defendida pelo professor, tem-se o pensamento crítico como algo que advém da capacidade humana de indignação, buscando compreendê-la e tornando a questioná-la numerosas vezes com o intuito de, justamente, buscar saber das tensões de sociedades e de sistemas usuais. Flores afirma que criticar não consiste em criar ou negar algo, mas ser criativo e afirmativo:

[...] Portanto, ser crítico exige afirmar os próprios valores como algo necessário a implementar lutas e garantias com todos os meios possíveis e, paralelamente, mostrar as contradições e as fraquezas dos argumentos e as práticas que a nós se opõem. É preciso afirmar as fraquezas de uma ideia, de um argumento, de um raciocínio, inclusive dos nossos, quando não forem consistentes, tentando corrigi-los para reforçá-los. Isso, porém, não nos deve dirigir unicamente à destruição daquilo que não nos convém como resultado de uma paixão cega, mas à prática de uma ação racional necessária para podermos avançar.²¹⁵

Referente aos cinco deveres básicos, o pesquisador informa que estes devem se unir às quatro condições recém citadas, a fim de chegar em “práticas emancipadoras baseadas nas lutas pela dignidade”, criando por si só zonas emancipadoras nas quais grupos sociais possam se sentir à vontade para buscar a igualdade de forma cultural, jurídica, social e política.²¹⁶

Os deveres abarcados na concepção em questão são o reconhecimento de que todos ajam culturalmente em relação ao mundo em que vivem; o respeito para inserir a última ideia de reconhecimento dos demais nas sociedades e práticas emancipadoras, identificando quem fica em eventuais posições de privilégio ou posições de subordinação de acordo com o sistema econômico-político vigente; a reciprocidade como forma de saber devolver o que acabou sendo de alguém por privilégio de maneira a contribuir com uma convivência mais realista e harmônica; a responsabilidade de um grupo assumir a opressão em relação a outro e, da mesma maneira, a responsabilidade deste outro de cobrar responsabilidades daqueles que cometeram “destruição as condições de vida dos demais”; e, por fim, a redistribuição de elementos práticos de garantia da dignidade humana e da igualdade social perante “regras jurídicas, fórmulas institucionais e ações políticas e econômicas concretas” que viabilizem o acesso aos bens hoje reconhecidos como essenciais, sobretudo dos não-essenciais que garantem o bem-estar e a dignidade humana.²¹⁷

²¹⁵ FLORES, Joaquín Herrera. **A (re)invenção dos direitos humanos**. Florianópolis: Fundação Boiteux, 2009, p. 66.

²¹⁶ FLORES, Joaquín Herrera. **A (re)invenção dos direitos humanos**. Florianópolis: Fundação Boiteux, 2009, p. 67.

²¹⁷ FLORES, Joaquín Herrera. **A (re)invenção dos direitos humanos**. Florianópolis: Fundação Boiteux, 2009, p. 68.

Através da profunda pesquisa dos direitos humanos por meio de um pensamento crítico, Flores criou uma ferramenta didática e visual a fim de compreender a complexidade e as várias camadas do tema pesquisado, aplicando-o na teoria e na prática. Chamado esse “esquema” de conhecimento e ação de diamante ético, destaca:

Com o “diamante ético”, nos lançamos a uma aposta: os direitos humanos vistos em sua real complexidade constituem um marco para construir uma ética que tenha como horizonte a consecução das condições para que “todos e todas” (indivíduos, culturas, formas de vida) possam levar à prática sua concepção de dignidade humana.²¹⁸

O esquema em questão possui três camadas distintas, mas com pontos de conexão entre elas. Em outras palavras, o pesquisador afirma que é uma figura tridimensional que está em constante movimento cuja representação é feita pelo “resultado de lutas que se sobrepõem com o passar do tempo”.²¹⁹

O resultado das diversas lutas em busca da emancipação humana é representado tanto pelos elementos teóricos considerados “conceituais” (linha vertical da imagem acima) quanto pelos elementos práticos nomeados “materiais” (linha horizontal da mesma imagem). No centro do diamante, mantém-se a ideia da dignidade humana, a qual retrata a conquista pela igualdade no que tange aos bens imateriais e materiais nas distintas maneiras de vida e de viver que existem.²²⁰ O autor complementa sobre:

O objetivo, portanto, dessa imagem metodológica se baseia na ideia de que tanto a dignidade humana como os direitos não são elementos isolados, e também, não são dados com antecedência, mas sim construídos passo a passo pela própria comunidade ou grupo afetado, o que lhes outorga um caráter de direitos em movimento que se podem gerar e revisar através da metodologia que se propõe.²²¹

A obra em questão oferece três modos de análise dos direitos humanos através do pensamento crítico que podem ser possíveis por meio do diamante ético, os quais serão utilizados para conduzir os próximos capítulos do presente trabalho:

- 1) Escolhendo relações concretas entre diferentes elementos (por exemplo, ideias, valores, práticas sociais);
- 2) Estudando camadas inteiras (a posição, a disposição, a narração e a historicidade de um determinado direito ou prática social);

²¹⁸ FLORES, Joaquín Herrera. **A (re)invenção dos direitos humanos**. Florianópolis: Fundação Boiteux, 2009, p. 119.

²¹⁹ FLORES, Joaquín Herrera. **A (re)invenção dos direitos humanos**. Florianópolis: Fundação Boiteux, 2009, p. 120.

²²⁰ FLORES, Joaquín Herrera. **A (re)invenção dos direitos humanos**. Florianópolis: Fundação Boiteux, 2009, p. 122.

²²¹ FLORES, Joaquín Herrera. **A (re)invenção dos direitos humanos**. Florianópolis: Fundação Boiteux, 2009, p. 123.

3) Entrecruzando diferentes camadas do diamante (por exemplo, as relações entre determinadas políticas de desenvolvimento dos direitos e as relações sociais de produção que predominam em espaços ou instituições concretas).²²²

Outrossim, Flores finaliza o capítulo do diamante ético explicando, mais uma vez, o objetivo do emprego desta estrutura e onde ela pode nos levar:

O desafio, então, consiste em saber escolher os elementos e os eixos que mais interessam ao trabalhar de forma analítica e prática os direitos humanos, entendidos como processos que abrem (ou fecham) espaços de luta pela dignidade humana. Escolhamos outros exemplos, apliquemos o diamante e, coletivamente, verifiquemos se existe ou não o suficiente grau de dignidade, núcleo para o qual convergem os diferentes elementos que compõem o “eixo conceitual” e o “eixo material” de nosso diamante.²²³

Isto posto, é importante mencionar que este viés da teoria crítica dos Direitos Humanos mostra-se de suma relevância para a construção da Inteligência Artificial inserida no âmbito jurídico, isto com o intuito de analisar as diversas situações que ocorrem no judiciário. E com isto, a percepção de Direitos Humanos base para o aprendizado de máquinas deve reconhecer as diferenças culturais de cada nação em que estará inserido, e por isto, compreende-se que a teoria crítica traz esse olhar fundado na realidade social.

Além disto, a ética é outro fator de grande relevância para a construção de uma IA de confiança, o que se abordará no tópico a seguir.

3.2 ÉTICA NO ÂMBITO JURÍDICO

Tratar sobre a compreensão da ética no âmbito jurídico é parte essencial para o reconhecimento de uma Inteligência Artificial de confiança, portanto, revela-se imperioso discutir a relação preconizada entre direito e moral.

No campo da deontologia²²⁴, percebe-se que as regras jurídicas não se apresentam de maneira estanque, isto é, do dever ser social, haja vista que existem outros discursos resultantes de práticas que ordenam ou indicam o comportamento humano. Tem-se, assim, a religião, a moral de um determinado grupo, as exigências de condutas no trabalho para obtenção de organização funcional, etc. Dentre esses

²²² FLORES, Joaquín Herrera. **A (re)invenção dos direitos humanos**. Florianópolis: Fundação Boiteux, 2009, p. 145.

²²³ FLORES, Joaquín Herrera. **A (re)invenção dos direitos humanos**. Florianópolis: Fundação Boiteux, 2009, p. 149.

²²⁴ Deontologia é uma filosofia que faz parte da filosofia moral contemporânea, que significa ciência do dever e da obrigação. É um tratado dos deveres e da moral. Uma teoria sobre as escolhas dos indivíduos, o que é moralmente necessário e serve para nortear o que realmente deve ser feito.

exemplos, desponta com elevado destaque o discurso jurídico-normativo, que está fundamentado na criação de mecanismos de subjetivação dos indivíduos, agrupando-se à ordem das regras chamada imperativas, decretadas de forma política e objetivamente apresentadas para que todos possam cumpri-las, salvo quando admitida excusa pela lei.²²⁵

No decorrer da história do pensamento humano, precipuamente sob o prisma da civilização ocidental, o paradigma ético sempre esteve intimamente associado à ideia de justiça, havendo, nesse contexto, aproximações e distanciamentos entre ambos os conceitos. O conceito de justiça, por sua vez, é desenvolvido de modo mais aprofundado pelo direito. É no âmbito jurídico, desse modo, que afloraram as maiores preocupações com a concretização da justiça.

A concepção de justiça reflete, no curso da história, uma extensa gama de expectativas delineadas ao longo do tempo, resultando em inúmeras correntes que buscam refletir sobre o justo e o injusto, dentre elas: teoria sofista, socrática, platônica, aristotélica, cristã, agostiniana, tomista, rousseauniana, kantiana, hegeliana, kelseniana, rawlsiana. Entretanto, essas doutrinas sofreram uma forte influência do pensamento hegemônico ocidental, destacando Gusmão²²⁶, em síntese, o seguinte esquema: a) de Platão advém uma herança segundo a qual a justiça é virtude suprema; b) de Aristóteles advém uma herança segundo a qual a justiça é igualdade/proporcionalidade; c) dos juristas romanos advém uma herança segundo a qual a justiça é vontade de dar a cada um o seu (*iustitia est constans et perpetua voluntas ius suum cuique tribuendi*).

Portanto, qualquer que seja o referencial que se adote, a questão inicial e central remete ao que entendemos por Justiça.

Em Aristóteles²²⁷, se tem um amplo desenvolvimento filosófico da busca pelo justo. Com efeito, esse filósofo se concentrou com afinco sobre o ideal de justiça e foi responsável por sistematizar um estudo aprofundado sobre o tema. A justiça poderia, então, ser definida por ele como o necessário equilíbrio de condições entre as pessoas; mas é também a garantia de participação na distribuição do poder entre os

²²⁵ BITTAR, Eduardo Carlos Bianca; ALMEIDA, Guilherme Assis de. **Curso de Filosofia do Direito**. 8. ed. rev. aum., São Paulo: Atlas, 2010, p. 439.

²²⁶ GUSMÃO, Paulo Dourado. **Introdução ao estudo do direito**. Rio de Janeiro: Forense, 1999, p. 71-72.

²²⁷ ARISTÓTELES. **Ética a Nicômaco**. Livro II. Coleção Os Pensadores, São Paulo: Nova Cultural, 1996.

indivíduos.

O justo nesta acepção é, portanto, o proporcional, e o injusto é o que viola a proporcionalidade. Neste último caso, um quinhão se torna muito grande e outro muito pequeno, como realmente acontece na prática, pois a pessoa que age injustamente fica com um quinhão muito grande do que é bom e a pessoa que é tratada injustamente fica com um quinhão muito pequeno. No caso do mal o inverso é verdadeiro, pois o mal maior, já que o mal menor deve ser escolhido em preferência ao maior, e o que é digno de escolha é um bem, e o que é mais digno de escolha é um bem ainda maior.²²⁸

Justiça, portanto, é antes de tudo a garantia da possibilidade de convivência entre desiguais, a máxima do senso comum - “alguns são mais iguais que outros”, tentando equiparar e padronizar respostas a problemas que se apresentam, forçosamente, na vida em sociedade.

A noção de justiça, porém, não consiste e nem deve consistir em dar a cada um o que é seu, do velho e superado brocardo romano – *suum cuique tribuere*. A longa experiência dos tempos mostra que a distribuição dos bens na sociedade não é equitativa, mas desigual e injusta, obediente a fatores diferenciados de poder, e não de critérios racionais e de contenção comum. Portanto, se prevalecesse a regra de dar a cada um o que é seu, sobretudo na sociedade exageradamente desigual de nossa época, consagrar-se-ia o reino da força e da esperteza. A justiça residiria em dar ao rico a sua fortuna e ao pobre a sua miséria, resultantes de desequilíbrios humanos moralmente insustentáveis.²²⁹

Para Bittar e Almeida o justo e o injusto, não se originam na natureza das coisas, mas nas opiniões e convenções do homem, ou seja, é definida por critérios humanos, e não naturais, pois se fossem naturais, todas as leis seriam iguais. Nesse sentido:

Em outras palavras, a mesma inconstância da legalidade (o que é lei hoje poderá não ser amanhã) passa a ser aplicada à justiça (o que é justo hoje poderá não ser amanhã). Nada do que se pode dizer absoluto (imutável, perene, eterno, incontestável...) é aceito pela sofística. Está aberto campo para o relativismo da justiça.²³⁰

Em Kelsen²³¹, no entanto, encontramos entendimento oposto. Para o autor austríaco, todo juízo de valor é irracional, posto que está diretamente embasado na fé. Dessa forma, não seria possível indicar cientificamente – ou seja, do modo racional - um valor que seja preterido em razão do outro.

Se no problema da justiça partirmos de um ponto de vista racional-científico,

²²⁸ ARISTÓTELES. **Ética a Nicômaco**. Livro II. Coleção Os Pensadores, São Paulo: Nova Cultural, 1996, p. 199.

²²⁹ PINHEIRO, Pe. José Ernanne Pinheiro; SOUSA JUNIOR, José Geraldo de; DINIS, Melillo et al. (Org.). **Ética, justiça e direito: reflexões sobre a reforma do judiciário**. 2. ed., Petrópolis: Vozes; CNBB, 1996, p. 128.

²³⁰ BITTAR, Eduardo Carlos Bianca; ALMEIDA, Guilherme Assis de. **Curso de Filosofia do Direito**. 8. ed. rev. aum., São Paulo: Atlas, 2010, p. 96.

²³¹ KELSEN, Hans. **O que é justiça?: a justiça, o direito e a política no espelho da ciência**. Tradução de Luís Carlos Borges. São Paulo: Martins Fontes, 1998.

não-metafísico, e reconhecermos que há muitos ideais de justiça diferentes uns dos outros e contraditórios entre si, nenhum dos quais exclui a possibilidade de um outro, então apenas nos será lícito conferir uma validade relativa aos valores da justiça constituídos através desses ideais.²³²

Sobre o assunto, Comparato destaca a existência de uma conexão entre os valores da justiça, da verdade e do amor:

Entre esses valores e princípios éticos não há concorrência, mas complementariedade. A justiça tende a se estiolar e, portanto, perder sua efetiva vigência se não for incessantemente aprofundada pelo amor. Este, por sua vez, descamba para um egoísmo disfarçado, ou um tíbio sentimentalismo, se não se fundar nas exigências primárias de justiça, das quais representa um aperfeiçoamento e jamais um sucedâneo. Como salientou Mahatma Gandhi, a ahimsa ou não-violência nada mais é do que o amor, entendido como um estado positivo de fazer o bem aos que nos ofendem ou prejudicam. Nessa concepção, a satyagraha, como disposição interior de amor incondicional à verdade, exige de todos os que a ela aderem uma ação incessante contra a injustiça, em qualquer de suas modalidades.²³³

Adverte Bittar e Almeida, contudo, ante a íntima relação entre direito e justiça, que:

A justiça não é coercível, é autônoma, correspondendo a uma norma moral, e não a uma norma jurídica. Normas jurídicas absorvem conteúdos de normas de justiça, funcionam como fonte de compelir coercitivamente comportamentos injustos, de proscrevê-los socialmente, mas não há que se negar a natureza da justiça como norma moral, e não jurídica.²³⁴

A concepção de justiça deve ser objeto busca incessante por parte dos operadores do sistema jurídico. Importa notar, todavia, que realmente pertinente é construir um modo do qual se possa realizá-la, principalmente, tendo como instrumento o direito. Gusmão²³⁵ resume a questão com bastante propriedade, sustentando que “o direito é norma executável coercitivamente, enquanto a justiça é um ideal, ou melhor, uma experiência constante, um valor, que pode ou não ser acolhido pelo legislador, apesar de dever sê-lo”.

Racionalizando e tentando padronizar o comportamento humano, a ética pretende efetuar uma crítica da Moral e, em certos casos, acaba por avançar no campo de atuação do direito. Assim:

O saber ético estuda o agir humano. Isso já se disse. Também já se disse que as normas morais convivem com as normas sociais. Porém, o que ainda está por ser dito, é que dentre as normas sociais e as demais convenções se

²³² Kelsen, Hans. **O que é justiça?: a justiça, o direito e a política no espelho da ciência**. Tradução de Luís Carlos Borges. São Paulo: Martins Fontes, 1998, p. 17.

²³³ Comparato, Fábio Konder. **Ética, Direito, Moral e Religião no Mundo Moderno**. São Paulo: Companhia das Letras, 2006, p. 521.

²³⁴ Bittar, Eduardo Carlos Bianca; Almeida, Guilherme Assis de. **Curso de Filosofia do Direito**. 8. ed. rev. aum., São Paulo: Atlas, 2010, p. 448.

²³⁵ Gusmão, Paulo Dourado. **Introdução ao estudo do direito**. Rio de Janeiro: Forense, 1999, p. 71.

destacam as normas jurídicas, com as quais interagem as normas morais.²³⁶

Para Reale²³⁷, a normativa moral não é assinalada por nenhum fenômeno exterior, não podendo ser sancionada ou promulgada, bem como não é cogente, o que inviabiliza, desde já, a punição por qualquer autoridade pública. De todo modo, a afronta a preceitos morais acaba sendo penalizada por outros meios de controle social - tais como rejeição, vergonha, exclusão de um determinado grupo, entre outras - que, em alguns casos, são até mesmo mais eficazes que a correção por via de normas jurídicas. Nesse sentido:

Sendo o direito um bem cultural, nele há sempre uma exigência axiológica atualizando-se na condicionalidade histórica, de maneira que a objetividade do vínculo jurídico está sempre ligada às circunstâncias de cada sociedade, aos processos de opção ou de preferência entre os múltiplos caminhos que, como vimos, se entreabrem no momento de qualquer realização de valores. Põe-se, assim, no âmago da experiência jurídica a problemática do Poder, que procura assegurar por todos os modos, pela força física, a realização do Direito.²³⁸

Reale²³⁹, sobre o assunto ainda acrescenta: “o direito, como experiência humana, situa-se no plano da ética, referindo-se a toda a problemática da conduta humana subordinada a normas de caráter obrigatório”.

Decorre disso que, do ponto de vista sociológico, o direito, como apreciação do fato ou fenômeno social, está subordinado à Moral, bem como à imposição de parâmetros nem sempre perfeitamente racionalizados, os quais se revelam vinculados a pressupostos de senso comum, míticos, religiosos, políticos, etc., conforme ensina Bittar:

O direito deve possuir como atributo constante um compromisso com a ética do coletivo. Isso significa que em suas estruturas devem estar as principais inquietações e principais premências gerais da sociedade; aos anseios sociais deve o direito responder como adequada, completa e eficaz normalização. Ou seja, está-se a discutir o compromisso que coloca o direito na frente de batalha pelos valores sociais mais caros a todos (saúde, educação, alimentação, higiene, saneamento, habitação, dignidade...), e não a um grupo, e não em favor de privilegiados, e não em detrimento de garantias fundamentais... É nessa ética do coletivo que os atos, as decisões, os entendimentos, as interpretações... devem se fiar no sentido da realização da tectura finalística, porém não idealista, e sim diária, do instrumental jurídico.²⁴⁰

A implicação recíproca entre direito e moral é destacada pelo autor Ferraz

²³⁶ BITTAR, Eduardo Carlos Bianca. **Curso de ética jurídica: ética geral e profissional**. 5. ed. rev., São Paulo: Saraiva, 2007, p. 32.

²³⁷ REALE, Miguel. **Filosofia do direito**. 19. ed., São Paulo: Saraiva, 1999, p. 703.

²³⁸ REALE, Miguel. **Filosofia do direito**. 19. ed., São Paulo: Saraiva, 1999, p. 703.

²³⁹ REALE, Miguel. **Filosofia do direito**. 19. ed., São Paulo: Saraiva, 1999, p. 37.

²⁴⁰ BITTAR, Eduardo Carlos Bianca. **Curso de ética jurídica: ética geral e profissional**. 5. ed. rev., São Paulo: Saraiva, 2007, p. 63.

Junior²⁴¹, que sustenta que o âmbito jurídico-normativo pode acompanhar os ditames morais de determinada sociedade, bem como pode deixar de observá-los, sendo que na primeira hipótese, tem-se um direito moral, inovador e democrático, e, na segunda, um direito imoral.

O referido autor complementa sua exposição afirmando que o direito imoral, ainda que venha a contestar os valores sedimentados na sociedade, é exigível e deve ser necessariamente cumprido, no intuito de impedir eventual sanção pela inobservância dos mandamentos jurídicos. Dessa forma, os direitos moral e imoral, sob esse prisma, são compreendidos como válidos, muito embora apresentem sentidos distintos, posto que o direito moral possui razão de existência, sendo ratificado e aceito pela sociedade, consubstanciando elemento instrumentalizador da justiça. Por outro lado, o direito imoral faz presumir que o ordenamento jurídico é uma via de poder e autoridade, que na realidade não detém legitimidade e é praticado mediante o uso da força, sem qualquer disposição prudencial.²⁴²

Destarte, o direito busca subsídios na moral. Esta, conforme se modifica, altera a própria legislação e a interpretação das leis. Pode-se dizer, nesse contexto, que direito sem moral resulta em puro arbítrio, e não direito, tanto que, como ciência humana, o direito é positivista, circunscrito ao poder de coerção que resulta da força das leis, dos costumes institucionalizados.

3.3 PARÂMETROS PARA REGULAMENTAÇÃO BRASILEIRA SOBRE O USO DA INTELIGÊNCIA ALRTIFICIAL COM BASE NAS ORIENTAÇÕES ÉTICAS DA COMISSÃO EUROPÉIA

Conforme já abordado no segundo capítulo, atualmente vem ganhando destaque o Projeto de Lei n. 21/2020, aprovado pela Câmara dos Deputados no dia 29 de setembro de 2021 e encaminhado para o Senado Federal para votação.

Ainda, cabe mencionar que em abril de 2021, o Ministério da Ciência, Tecnologia e Inovações, vinculado ao Governo Federal, apresentou a Estratégia Brasileira de Inteligência Artificial (EBIA)²⁴³ com o objetivo de traçar um plano de

²⁴¹ FERRAZ JUNIOR, Tércio Sampaio. **Introdução ao estudo do direito**: técnica, decisão, dominação. 2. ed., São Paulo: Atlas, 1994, p. 326.

²⁴² FERRAZ JUNIOR, Tércio Sampaio. **Introdução ao estudo do direito**: técnica, decisão, dominação. 2. ed., São Paulo: Atlas, 1994, p. 326.

²⁴³ BRASIL. Ministério da Ciência, Tecnologia e Inovações. **Estratégia Brasileira de Inteligência**

desenvolvimento do país neste âmbito. O documento norteará a atuação do poder executivo brasileiro no que diz respeito ao desenvolvimento de ações que estimulem a pesquisa, a inovação e o desenvolvimento de soluções em IA, bem como o uso consciente, ético e em prol de um futuro melhor. A EBIA foi construída colhendo visões diversas e setoriais, inclusive considerando experiências internacionais.

Ressalta-se que a EBIA tem como objetivos: contribuir para a elaboração de princípios éticos para o desenvolvimento e uso de IA responsável; promover investimentos sustentados em pesquisa e desenvolvimento em IA; remover barreiras à inovação em IA; capacitar e formar profissionais para o ecossistema da IA; estimular a inovação e o desenvolvimento da IA brasileira em ambiente internacional; e promover ambiente de cooperação entre os entes públicos e privados, a indústria e os centros de pesquisas para o desenvolvimento da IA. No documento, são estabelecidos 9 eixos temáticos e um conjunto de 74 ações estratégicas. Além da EBIA, estão sendo criados 8 centros de IA no Brasil.²⁴⁴

No entanto, para que haja um desenvolvimento da IA que vise o bem comum em prol da humanidade deve, sem dúvida, considerar as necessidades humanas. Em virtude disso, com intuito de maximizar os benefícios da IA e, simultaneamente, prevenir e minimizar seus riscos, surge a necessidade de criar um sistema de normas que guiam o projeto de uma IA confiável.

Sendo assim, a construção de uma Inteligência Artificial deve essencialmente garantir o respeito aos Direitos Humanos, bem como a garantia dos Direitos Fundamentais previstos na Constituição Federal Brasileira de 1988. Com base nesta perspectiva observa-se que a Ética é caminho essencial para a construção de uma IA de confiança (termo utilizado pela Comissão Europeia ao criar o documento denominado “Orientações Éticas para uma IA de Confiança”).

E ao tratar da compreensão dos Direitos Humanos, não se pode esquecer que em um mundo globalizado, a IA serve para todo tipo de cultura social, para todo tipo de governo, inclusive para aqueles que não se estruturam em valores relacionados a direitos humanos, o mundo virtual não possui fronteiras. Por isso, se deu relevância a

Artificial (EBIA). 2021. Disponível em: https://www.gov.br/mcti/pt-br/acompanhe-o-mcti/transformacaodigital/arquivosinteligenciaartificial/ia_estrategia_documento_referencia_4-979_2021.pdf. Acesso em: 09 jan. 2022.

²⁴⁴ BRASIL. Ministério da Ciência, Tecnologia e Inovações. **Estratégia Brasileira de Inteligência Artificial (EBIA).** 2021. Disponível em: https://www.gov.br/mcti/pt-br/acompanhe-o-mcti/transformacaodigital/arquivosinteligenciaartificial/ia_estrategia_documento_referencia_4-979_2021.pdf. Acesso em: 09 jan. 2022.

teoria crítica dos Direitos Humanos de Joaquín Herrera Flores, por ser desenvolvida com intenção de romper com o pensamento hegemônico dos direitos universais fundamentados nas noções ocidentais e eurocêtricas de dignidade humana, que alicerça o sistema neoliberal intensificando as desigualdades sociais. Desta forma, ao se pensar em uma regulamentação para a IA, mesmo que seja uma regulamentação interna, se faz necessário ter essa compreensão global.

Isso é relevante à medida que o sistema hegemônico da atualidade é, em sua grande parte, neoliberal. Consequentemente, se dá menor importância para políticas públicas de igualdade social, econômica e cultural, o que deixa as normas jurídicas sem efetiva aplicação, e isso se amplia quando consideramos a abrangência do ambiente virtual. As normas determinam que o acesso aos meios jurídicos deveria ser igual para todos, mas, na realidade, nem todos têm “por igual os direitos, ou seja, os instrumentos e meios para levar adiante nossas lutas pelo acesso aos bens necessários para afirmar nossa própria dignidade”. É dessa lógica que deriva a realidade de certos grupos, que enfrentam maiores dificuldades de utilização dos meios jurídicos para garantirem acesso aos bens exigíveis para uma vida digna.²⁴⁵

O uso da IA deve ser harmônico com os direitos garantidos aos cidadãos, sejam aqueles consagrados na constituição nacional ou pela Declaração Universal de Direitos Humanos. A utilização de novas tecnologias pelo mercado ou pelos Estados devem sempre estar alinhadas com o bem-estar social e com a promoção de uma vida digna, para isto, é necessário que se construa uma legislação bom base em parâmetros de uma IA de Confiança, tendo em vista que a confiança e a ética são os pontos centrais para a aplicação da inteligência artificial.

Foi neste contexto que a União Europeia (GPAN IA – Comissão Europeia) juntou um grupo de peritos de alto nível para discutirem e chegaram a um comum acordo sobre a aplicação da Inteligência Artificial. E é com base neste artigo oriundo do “Grupo Independente de Especialistas de Alto Nível em Inteligência Artificial”, nomeado pela Comissão Europeia, que se pode concluir a necessidade de uma confiabilidade entre o ser humano, tanto aquele que manuseia a máquina, como aquele que é diretamente ou indiretamente afetado por ela e a própria inteligência artificial.²⁴⁶

²⁴⁵ FLORES, Joaquín Herrera. **A (re)invenção dos direitos humanos**. Florianópolis: Fundação Boiteux, 2009, p. 41.

²⁴⁶ COMISSÃO EUROPEIA. **Orientações éticas para uma IA de Confiança**. Junho de 2018. Disponível em: <http://escola.mpu.mp.br/servicos-academicos/atividades->

Segundo o documento, os sistemas de IA devem ser “robustos” e “seguros”, de modo a evitar erros ou a terem condição de lidar com estes, corrigindo eventuais inconsistências. Esses problemas podem ter sérios impactos na sociedade, como a discriminação de pessoas no acesso a um serviço ou até mesmo quedas de bolsas cujas compras e vendas de ações utilizam essas tecnologias.²⁴⁷

Não obstante a isso, o documento traz uma perspectiva de que: para que esta confiança seja assente, terá que se fundar no interesse público e nos valores destes indivíduos conforme suas garantias e direitos fundamentais, tendo como norte a prevalência dos Direitos Humanos.

3.3.1 Componentes de uma Inteligência Artificial de confiança

Diante desse contexto, uma IA de confiança tem três componentes, que devem ser observados ao longo de todo o ciclo de vida do sistema:

1. Deve ser legal, garantindo o respeito de toda a legislação e regulamentação aplicáveis;
2. Deve ser ética, garantindo a observância de princípios e valores éticos; e
3. Deve ser sólida, tanto do ponto de vista técnico como do ponto de vista social, uma vez que, mesmo com boas intenções, os sistemas de IA podem causar danos não intencionais.²⁴⁸

Cada um destes três componentes é necessário para alcançar uma IA de confiança, porém, não são suficientes isoladamente, o ideal é que as três funcionem em harmonia e se sobreponham durante sua ação. Estes três componentes formam a base para a IA de confiança, que devem ser pautados nos direitos fundamentais, no Brasil, consagrados na Constituição Federal de 1988 e nos Direitos Humanos. É evidente que podem surgir tensões entre tais componentes, tendo-se a ponderação como técnica para alinhá-los.

Sendo assim, a IA legal trata da garantia ao respeito a legislação e regulamentação aplicável. Hoje, no cenário brasileiro, como já abordado no capítulo 2, a regulamentação sobre a Inteligência Artificial é escassa, todavia, há de se levar

academicas/inovaescola/curadoria/3-ciclo-de-debates/inteligencia-artificial-e-internet-das-coisas-oportunidades-e-desafios/ethicsguidelinesfortrustworthyai-ptpdf.pdf. Acesso em: 20 out. 2021.

²⁴⁷ VALENTE, Jonas. **Europa lança diretrizes éticas para o uso da inteligência artificial**. 2019. Disponível em: <https://agenciabrasil.ebc.com.br/internacional/noticia/2019-04/europa-lanca-diretrizes-eticas-para-o-uso-da-inteligencia-artificial>. Acesso em: 10 nov. 2021.

²⁴⁸ COMISSÃO EUROPEIA. **Orientações éticas para uma IA de Confiança**. Junho de 2018. Disponível em: <http://escola.mpu.mp.br/servicos-academicos/atividades-academicas/inovaescola/curadoria/3-ciclo-de-debates/inteligencia-artificial-e-internet-das-coisas-oportunidades-e-desafios/ethicsguidelinesfortrustworthyai-ptpdf.pdf>. Acesso em: 20 out. 2021, p. 06.

em consideração o que determina a Constituição Federal em seu artigo 5º acerca dos direitos e garantias fundamentais:

Art. 5º Todos são iguais perante a lei, sem distinção de qualquer natureza, garantindo-se aos brasileiros e aos estrangeiros residentes no País a inviolabilidade do direito à vida, à liberdade, à igualdade, à segurança e à propriedade, nos termos seguintes:²⁴⁹

Verifica-se que é imprescindível que não haja violação às disposições do art. 5º da Constituição Federal, além daquelas decorrentes do regime e dos princípios adotados pela Constituição Federal de 1988. Não se ignora que podem surgir tensões em tal processo, especialmente em se tratando de tipicidade aberta, o que demanda do intérprete um esforço argumentativo que compatibilize os interesses em questão. Nesse contexto, elenca-se inicialmente a necessidade de respeito à dignidade da pessoa humana, enquanto sujeito que possui valor intrínseco e que não pode ser reprimido por tecnologias de inteligência artificial.²⁵⁰

Uma das grandes preocupações que existe na utilização da IA decorre da possibilidade de responsabilização acerca da conduta realizada. No instante em que as funções de uma IA se equiparam a uma conduta humana, estas logicamente, demonstram-se capazes de realizar as mesmas atividades, sendo assim, podendo incorrer em algum momento nos mesmos equívocos, a diferença que uma IA, independente da sua forma, não poderá ser culpabilizada por estes erros, o que significa que as consequências ainda não podem ser dimensionadas quanto a transgressão de direitos fundamentais das partes envolvidas.²⁵¹

Também se ressalta a necessidade de respeito pelo regime democrático, de modo que os sistemas de Inteligência artificial devem respeitar a pluralidade de valores e escolhas individuais, sem prejudicar compromissos fundamentais sobre os quais o Estado Democrático de Direitos se fundamenta, tais como a isonomia e o devido processo legal.²⁵²

²⁴⁹ BRASIL. [Constituição, 1988]. **Constituição da República Federativa do Brasil de 1988**. Disponível em: http://www.planalto.gov.br/ccivil_03/constituicao/constituicao.htm. Acesso em: 20 out. 2021.

²⁵⁰ PEREIRA, Thiago Pedroso. **A legalidade e efetividade dos atos judiciais realizados por inteligência artificial**. 2020. 122 fls. Dissertação (Mestrado em Direito) – Universidade Nove de Julho, São Paulo, 2020. Disponível em: <https://bibliotecatede.uninove.br/bitstream/tede/2411/2/Thiago%20Pedroso%20Pereira.pdf>. Acesso em: 27 dez. 2021.

²⁵¹ SILVA, Gabriela Buarque Pereira; EHRHARDT JÚNIOR, Marcos. Diretrizes éticas para a inteligência artificial confiável na União Europeia e a regulação jurídica no Brasil. *In: Revista IBERC*, v. 3, n. 3, p. 1-28, 2020. Disponível em: <https://revistaiberc.responsabilidadecivil.org/iberc/article/view/133/105>. Acesso em: 27 dez. 2021.

²⁵² PEREIRA, Thiago Pedroso. **A legalidade e efetividade dos atos judiciais realizados por**

Tal preocupação é especialmente relevante quando se constata que a inteligência artificial vai muito além das recomendações que são apresentadas no Spotify ou na Netflix. Evidencia-se a crescente utilização da inteligência artificial por órgãos judiciais, como por exemplo, e já explicado neste trabalho, no Supremo Tribunal Federal desenvolveu-se um sistema chamado VICTOR, cujo objetivo, conforme já mencionado anteriormente, é ler os recursos extraordinários que chegam a essa Corte e identificar a vinculação com determinados temas de repercussão geral.

A legislação estabelece obrigações positivas e negativas, ou seja, deve ser interpretada não só à luz do que não pode ser feito, mas também do que deve ser feito. A legislação não só proíbe certas ações como também permite outras. Saliente-se, a este respeito, que a regulamentação da IA deve estar de acordo com a legislação vigente, procurando garantir a fiabilidade da IA, tais como a proteção de dados e a não discriminação.

No que se refere ao componente da IA ética, esta trata da garantia da observância dos princípios e valores éticos. Evidencia que a legislação não avança na mesma velocidade da evolução tecnológica, restando, por sua vez, muitas vezes desatualizada ou desadequada a realidade presenciada pela sociedade.

Para alcançar uma IA de confiança é necessário cumprir a legislação, mas não só, pois esta é apenas uma das suas três componentes. A legislação nem sempre acompanha a rapidez da evolução tecnológica e, por vezes, pode estar defasada em relação as normas éticas ou não se adequar, pura e simplesmente, ao tratamento de certas questões. Por conseguinte, para os sistemas de IA serem confiáveis, devem também ser éticos e estar em harmonia com normas éticas.²⁵³

Foi neste sentido, visando a garantia da aplicação de uma IA legal e ética que no ano de 2020, o CNJ publicou a Resolução n. 332²⁵⁴, que dispõe acerca da ética, da transparência e da governança na produção e no uso de Inteligência Artificial no

inteligência artificial. 2020. 122 fls. Dissertação (Mestrado em Direito) – Universidade Nove de Julho, São Paulo, 2020. Disponível em: <https://bibliotecatede.uninove.br/bitstream/tede/2411/2/Thiago%20Pedroso%20Pereira.pdf>. Acesso em: 27 dez. 2021.

²⁵³ COMISSÃO EUROPEIA. **Orientações éticas para uma IA de Confiança.** Junho de 2018. Disponível em: <http://escola.mpu.mp.br/servicos-academicos/atividades-academicas/inovaescola/curadoria/3-ciclo-de-debates/inteligencia-artificial-e-internet-das-coisas-oportunidades-e-desafios/ethicsguidelinesfortrustworthyai-ptpdf.pdf>. Acesso em: 20 out. 2021, p. 11.

²⁵⁴ RASIL. Conselho Nacional de Justiça. **Resolução nº 332 de 21 de agosto de 2020.** Dispõe sobre a ética, a transparência e a governança na produção e no uso de Inteligência Artificial no Poder Judiciário e dá outras providências. Disponível em: <https://atos.cnj.jus.br/atos/detalhar/3429>. Acesso em: 08 set. 2021.

Poder Judiciário e dá outras providências, todavia esta resolução, como é sabido não tem força de lei, visando apenas regular uma lei futura, o que ainda não ocorreu.

É certo afirmar que as IA's apresentam soluções eficientes e eficazes para os mais diversos problemas, contudo, o que demanda essa discussão é o fato que essa autonomia possa ensejar questões de ordem ética ainda não aventadas pela legislação.

A questão da ética é um dos pilares da IA de confiança, e a preocupação com a ética pode se ver em evidência na Portaria GM nº 4.617/2021, nas ações estratégicas estabelecidas:

- Estimular a produção de uma IA ética financiando projetos de pesquisa que visem a aplicar soluções éticas, principalmente nos campos de equidade/não-discriminação (fairness), responsabilidade/prestação de contas (accountability) e transparência (transparency), conhecidas como a matriz FAT.
- Estimular parcerias com corporações que estejam pesquisando soluções comerciais dessas tecnologias de IA ética.
- Estabelecer como requisito técnico em licitações que os proponentes ofereçam soluções compatíveis com a promoção de uma IA ética (por exemplo, estabelecer que soluções de tecnologia de reconhecimento facial adquiridas por órgãos públicos possuam um percentual de falso positivo abaixo de determinado limiar).²⁵⁵

No entanto, para que se possa garantir a ética é necessário que se institua uma comissão de ética, para analisar os processos de revisão ética, direcionada para a construção da IA de confiança, o que já se encontra estabelecido na Portaria GM nº 4.617/2021. Muitas organizações já têm ou estão considerando a criação de conselhos de revisão de dados ou comitês de ética em relação à IA, que podem ser internos ou externos a tais organizações. Essa é vista como uma maneira de impulsionar accountability dentro das corporações, promover tomadas de decisões responsáveis e garantir que novas utilizações de dados respeitem os valores corporativos e sociais.²⁵⁶

Embora o fato de ser uma tecnologia de utilização em diversos países atualmente, que automatizaram seus processos judiciais, é de suma importância salientar que estes mesmos, tem demonstrado grande preocupação com questões éticas advindas destes novos processos. Tendo em vista a interação que estes

²⁵⁵ BRASIL. **Portaria GM nº 4.617, de 6 de abril de 2021.** Institui a Estratégia Brasileira de Inteligência Artificial e seus eixos temáticos. Disponível em: https://www.in.gov.br/en/web/dou/-/portaria-gm-n-4.617-de-6-de-abril-de-2021-*-313212172. Acesso em: 27 dez. 2021.

²⁵⁶ BRASIL. **Portaria GM nº 4.617, de 6 de abril de 2021.** Institui a Estratégia Brasileira de Inteligência Artificial e seus eixos temáticos. Disponível em: https://www.in.gov.br/en/web/dou/-/portaria-gm-n-4.617-de-6-de-abril-de-2021-*-313212172. Acesso em: 27 dez. 2021.

sistemas, denominados IA's têm com a coleta, tratamento e análise dos dados utilizados para que seja realizado a tomada efetiva de decisão.²⁵⁷

Ainda que as IA's utilizadas para essa tarefa, comumente, possuam capacidade de aprendizado, ou, como dito anteriormente neste estudo, *machine learning*, por serem desenvolvidos sobre um algoritmo, tendenciam a apresentar resultados lógicos, distantes complexos e de difícil entendimento ao ser humano desprovido de conhecimento nesse sentido. Ainda que de forma mínima, existe a possibilidade também de que a base logarítmica onde a IA foi elaborada contenha erros, o que, mesmo que a máquina possa aprender, o fará com erros e/ou dificuldades pelo conflito em seus parâmetros base.

Desta forma, o que se destaca como maior necessidade de discussão, é que o judiciário demonstre a seus jurisdicionados, de forma clara e assertiva quais os parâmetros utilizados para se chegar a estes resultados, ou seja, demonstrar ao grande público e agentes de utilização como advogados, procuradores, defensores, etc., razões para a automação, além de outros como a possibilidade ou não de revisão de atos tomados unicamente por IA. Todos estes cuidados devem ter como única premissa, a de fazer com que a utilização de IA seja fonte de segurança a todos que se utilizam desse poder, e não mais um elemento de desconfiança, o que tornaria seu emprego inviável e obstaria os benefícios que o mesmo proporcionaria.

Assim, o terceiro componente ressalta que a IA deve ser sólida ou robusta, tanto para um ponto de vista técnico como para um ponto de vista social, uma vez que, mesmo sendo programada para obter resultados com boas intenções, os sistemas inteligentes estão sujeitos a ocasionarem danos não intencionais.²⁵⁸

Cabe ressaltar que estes são de suma importância, pois, além da necessária observância aos princípios éticos, é essencial que os sistemas de IA assegurem que seu funcionamento seja seguro e confiável, adequando-se à requisitos que previnam quaisquer impactos adversos não intencionais.

Segundo as Diretrizes Éticas da União Europeia:

²⁵⁷ PEREIRA, Thiago Pedroso. **A legalidade e efetividade dos atos judiciais realizados por inteligência artificial**. 2020. 122 fls. Dissertação (Mestrado em Direito) – Universidade Nove de Julho, São Paulo, 2020. Disponível em: <https://bibliotecatede.uninove.br/bitstream/tede/2411/2/Thiago%20Pedroso%20Pereira.pdf>. Acesso em: 27 dez. 2021.

²⁵⁸ COMISSÃO EUROPEIA. **Orientações éticas para uma IA de Confiança**. Junho de 2018. Disponível em: <http://escola.mpu.mp.br/servicos-academicos/atividades-academicas/inovaescola/curadoria/3-ciclo-de-debates/inteligencia-artificial-e-internet-das-coisas-oportunidades-e-desafios/ethicsguidelinesfortrustworthyai-ptpdf.pdf>. Acesso em: 20 out. 2021, p. 09.

2. Robustez técnica e segurança:

- 2.1. Avaliação de possíveis formas de ataque e vulnerabilidades do sistema de IA e as respectivas medidas para garantir a integridade e a resiliência do sistema contra tais disfunções, descrevendo riscos e seguranças;
- 2.2. Avaliação do comportamento do sistema em situações e ambientes inesperados;
- 2.3. Avaliação de plano de recuperação na hipótese de ataques adversários ou situações inesperadas;
- 2.4. Avaliação da comunicação ao usuário dos riscos apresentados pelo sistema de IA e da existência de planos para mitigar ou gerenciar tais riscos;
- 2.5. Avaliação da presença de apólices de seguro para lidar com possíveis danos do sistema de IA;
- 2.6. Avaliação da probabilidade de o sistema de IA causar danos aos usuários ou a terceiros, bem como ao meio ambiente ou aos animais;²⁵⁹

O requisito da robustez e segurança, por sua vez, inclui a proteção aos ataques de hackers, planos de retorno e confiabilidade. A robustez técnica está intimamente ligada ao princípio da prevenção de danos, porquanto determina que os sistemas possuam abordagem preventiva de riscos e que se comportem de modo confiável, minimizando a ocorrência de danos inesperados. É imprescindível, portanto, que se desenvolvam proteções contra vulnerabilidades, máxime considerando a amplitude e influência que o *hacking* pode ter sobre o funcionamento da máquina. Medidas de segurança insuficientes também podem ensejar decisões equivocadas ou até danos físicos ao usuário.

O nível de medidas de segurança requeridas em determinado sistema de IA deverá ser proporcional ao nível da magnitude do risco apresentado pela máquina. O risco será maior na medida em que a máquina tiver menor precisão em seu funcionamento, isto é, menor capacidade de fazer julgamentos e classificações corretas com base nos dados ou modelos. As Diretrizes estipulam, ainda, que quando previsões imprecisas ocasionais não puderem ser evitadas, é importante que o sistema possa indicar a probabilidade de ocorrência desses erros, o que concretiza a noção de boa-fé objetiva. É de extrema relevância que haja um mínimo controle sobre aquilo que se produz, para que seja possível, de modo transparente, informar os interessados sobre as limitações verificadas no produto ou serviço.²⁶⁰

Na mesma linha do documento europeu, a Portaria GM nº 4.617/2021 também

²⁵⁹ COMISSÃO EUROPEIA. **Orientações éticas para uma IA de Confiança**. Junho de 2018. Disponível em: <http://escola.mpu.mp.br/servicos-academicos/atividades-academicas/inovaescola/curadoria/3-ciclo-de-debates/inteligencia-artificial-e-internet-das-coisas-oportunidades-e-desafios/ethicsguidelinesfortrustworthyai-ptpdf.pdf>. Acesso em: 20 out. 2021, p. 09.

²⁶⁰ SILVA, Gabriela Buarque Pereira; EHRHARDT JÚNIOR, Marcos. Diretrizes éticas para a inteligência artificial confiável na União Europeia e a regulação jurídica no Brasil. *In: Revista IBERC*, v. 3, n. 3, p. 1-28, 2020. Disponível em: <https://revistaiberc.responsabilidadecivil.org/iberc/article/view/133/105>. Acesso em: 27 dez. 2021.

apresenta a robustez como essencial para que se possa ter confiança na implementação da IA: “Os sistemas de IA devem funcionar de maneira robusta, segura e protegida ao longo de seus ciclos de vida, e os riscos em potencial devem ser avaliados e gerenciados continuamente”.²⁶¹

Estes são apenas os componentes essenciais para a construção de uma IA de confiança, no entanto, é preciso estabelecer os princípios que serão a base da IA, assentes nos direitos fundamentais e os direitos humanos, que devem ser respeitados para garantir uma IA ética e sólida.

3.3.2 Princípios éticos bases de uma Inteligência Artificial de Confiança

Falar sobre a necessidade do uso ético da IA parece ser óbvio, no entanto, na realidade, é mais complexo integrar esses princípios éticos na tecnologia em si. As organizações reconhecem a importância de um enfoque holístico para gerenciar e governar as soluções de IA no ciclo de vida completo da IA.

Os princípios éticos estabelecidos no documento europeu formam a base essencial para a construção de uma IA que vise o bem-estar social e o desenvolvimento sustentável. Conforme descreve o documento:

A reflexão ética sobre a tecnologia de IA pode ter múltiplas finalidades. Em primeiro lugar, pode estimular a reflexão sobre a necessidade de proteger as pessoas e os grupos ao nível mais básico. Em segundo lugar, pode estimular novos tipos de inovações que procurem promover valores éticos, como os que contribuem para realizar os Objetivos de Desenvolvimento Sustentável da ONU13, que estão firmemente incorporados na próxima Agenda 2030 da UE14. Embora o presente documento se ocupe principalmente da primeira finalidade referida, a importância que a ética poderá ter na segunda finalidade não deve ser subestimada. Uma IA de confiança pode melhorar o desenvolvimento individual e o bem-estar coletivo mediante a geração de prosperidade, a criação de valor e a maximização da riqueza. Pode contribuir para alcançar uma sociedade justa, ajudando a aumentar a saúde e o bem-estar dos cidadãos de forma a fomentar a igualdade na distribuição das oportunidades económicas, sociais e políticas.²⁶²

O documento elaborado pela Comissão Europeia elenca como base para uma IA de confiança, quatro princípios éticos: o respeito da autonomia humana; a prevenção de dados; a equidade e a explicabilidade.

²⁶¹ BRASIL. **Portaria GM nº 4.617, de 6 de abril de 2021.** Institui a Estratégia Brasileira de Inteligência Artificial e seus eixos temáticos. Disponível em: <https://www.in.gov.br/en/web/dou/-/portaria-gm-n-4.617-de-6-de-abril-de-2021-313212172>. Acesso em: 27 dez. 2021.

²⁶² COMISSÃO EUROPEIA. **Orientações éticas para uma IA de Confiança.** Junho de 2018. Disponível em: <http://escola.mpu.mp.br/servicos-academicos/atividades-academicas/inovaescola/curadoria/3-ciclo-de-debates/inteligencia-artificial-e-internet-das-coisas-oportunidades-e-desafios/ethicsguidelinesfortrustworthyai-ptpdf.pdf>. Acesso em: 11 jan. 2022, p. 14.

É fundamental que haja um conhecimento mínimo acerca das capacidades e limitações da inteligência artificial, com o objetivo de facilitar a rastreabilidade e a auditabilidade dos sistemas de IA, especialmente em situações críticas. Não se ignora que muitas vezes não há conhecimento exato acerca de como essas máquinas funcionam, fenômeno chamado de black box da inteligência artificial. Em comentário a tal fenômeno, Will Knight argumenta que “Nós podemos construir esses modelos, mas nós não sabemos como eles trabalham”.²⁶³

A avaliação e a pesquisa nesse ramo, portanto, são essencialmente dinâmicas e visam continuamente a implementar requisitos e avaliar soluções, com o objetivo de que melhores resultados sejam alcançados para todas as partes envolvidas. As Diretrizes também ressaltam a necessidade de que os sistemas de inteligência artificial sejam centrados no ser humano, apoiados no compromisso de servir à humanidade e sua liberdade. Trata-se da perspectiva antropocêntrica que predomina no contexto pós-moderno, em que se alça a dignidade humana como epicentro do sistema normativo.

No que tange aos direitos fundamentais a serem observados, impende ressaltar que os tratados da União Europeia e a Carta da União Europeia prescrevem uma série de direitos fundamentais que os Estados-membros e as instituições da União Europeia são obrigados a respeitar, no tocante à dignidade, liberdade, igualdade e solidariedade. O alicerce de tais direitos reside precisamente na dignidade humana e os Direitos Humanos, que atribui ao ser humano uma posição única e inalienável, central na vida social, política e econômica.

No caso do Brasil, verifica-se que é imprescindível que não haja violação às disposições do art. 5º da Constituição Federal, bem como aos Direitos Humanos, além daquelas decorrentes do regime e dos princípios adotados pela Carta. Não se ignora que podem exsurgir tensões em tal processo, máxime em se tratando de tipicidade aberta, o que demanda do intérprete um esforço argumentativo que compatibilize os interesses em questão.

De acordo com Gabriela Buarque Pereira Silva e Marcos Ehrhardt Júnior, é fundamental que também haja uma postura de ceticismo acerca da concepção de neutralidade dos dados. Isso porque a inteligência artificial se baseia numa grande

²⁶³ KNIGHT, Will. The Dark Secret at the Heart of AI. *In: MIT Technology Review*, 2017. Disponível em: <https://www.technologyreview.com/2017/04/11/5113/the-dark-secret-at-the-heart-of-ai/>. Acesso em: 11 jan. 2022.

quantidade de dados e informações cuja mineração depende, sobretudo, de escolhas dos programadores. A operação depende essencialmente de inputs e de outputs do programador. Se os dados subjacentes são tendenciosos, as desigualdades estruturais e os preconceitos inculcados nos dados serão amplificados por meio da atividade da inteligência artificial. As próprias escolhas sobre inserção, organização e classificação de dados devem ser feitas de modo cauteloso por todos os envolvidos, sob pena de violação aos direitos de personalidade.²⁶⁴

Sendo assim, no que se refere ao princípio do respeito pela autonomia humana, “os seres humanos que interagem com os sistemas de IA devem ser capazes de manter uma autodeterminação plena e efetiva sobre si mesmos e poder participar do processo democrático”. Não poderia haver, portanto, subordinação ou manipulação dos seres humanos por meio da inteligência artificial, devendo esta servir para complementar e fomentar as habilidades cognitivas, sociais e culturais dos agentes, deixando margem de escolha ao ser humano.²⁶⁵

A distribuição de funções entre os seres humanos e os sistemas de IA devem seguir princípios de conexão centrados no ser humano e deixar uma oportunidade significativa para a escolha humana. Isto implica que se garanta a supervisão e o controle por parte de seres humanos sobre os processos de trabalho dos sistemas de IA.²⁶⁶

Atualmente as IA's disponíveis tem tamanho grau de autonomia que podem gerenciar tranquilamente uma carteira de processos, decidindo “*just in time*” as condutas a se realizar, inclusive de elaborar sentenças com base nos fatos e provas apresentadas.²⁶⁷

Sobre o princípio da prevenção de danos, de acordo com o documento europeu, os sistemas de IA não devem causar danos ou agravá-los nem afetar

²⁶⁴ SILVA, Gabriela Buarque Pereira; EHRHARDT JÚNIOR, Marcos. Diretrizes éticas para a inteligência artificial confiável na União Europeia e a regulação jurídica no Brasil. **In: Revista IBERC**, v. 3, n. 3, p. 1-28, 2020. Disponível em: <https://revistaiberc.responsabilidadecivil.org/iberc/article/view/133/105>. Acesso em: 11 jan. 2022.

²⁶⁵ COMISSÃO EUROPEIA. **Orientações éticas para uma IA de Confiança**. Junho de 2018. Disponível em: <http://escola.mpu.mp.br/servicos-academicos/atividades-academicas/inovaescola/curadoria/3-ciclo-de-debates/inteligencia-artificial-e-internet-das-coisas-oportunidades-e-desafios/ethicsguidelinesfortrustworthyai-ptpdf.pdf>. Acesso em: 11 jan. 2022, p. 17.

²⁶⁶ COMISSÃO EUROPEIA. **Orientações éticas para uma IA de Confiança**. Junho de 2018. Disponível em: <http://escola.mpu.mp.br/servicos-academicos/atividades-academicas/inovaescola/curadoria/3-ciclo-de-debates/inteligencia-artificial-e-internet-das-coisas-oportunidades-e-desafios/ethicsguidelinesfortrustworthyai-ptpdf.pdf>. Acesso em: 11 jan. 2022, p. 17.

²⁶⁷ PEREIRA, Thiago Pedroso. **A legalidade e efetividade dos atos judiciais realizados por inteligência artificial**. 2020. 122 fls. Dissertação (Mestrado em Direito) – Universidade Nove de Julho, São Paulo, 2020. Disponível em: <https://bibliotecatede.uninove.br/bitstream/tede/2411/2/Thiago%20Pedroso%20Pereira.pdf>. Acesso em: 11 jan. 2022.

negativamente os seres humanos de qualquer outra forma.²⁶⁸

A prevenção dos danos é um imperativo cada vez mais constante na contemporânea sociedade de risco. Diariamente surgem notícias acerca de ataques de *hackers* ou vazamentos indevidos de dados, o que seguramente tem o condão de violar direitos de personalidade dos usuários. Com efeito, sob a perspectiva de Ulrich Beck em sua obra “A sociedade de risco”, a sociedade contemporânea é marcada por perigos que se situam na imbricação entre construções científicas e sociais, sendo o desenvolvimento tecnológico uma fonte de causa, definição e solução de riscos. A IA deve assumir, nesse contexto, protagonismo na tentativa de mitigação e gerenciamento de crises.²⁶⁹

Portanto, o princípio da prevenção de danos baseia-se na consideração tanto do ambiente de implementação do sistema de IA quanto dos seres vivos que o envolvem.

Outro princípio presente na Carta de Ética da União Europeia é o princípio da justiça ou princípio da equidade, que significa dizer que o desenvolvimento, a implantação e a utilização dos sistemas de IA devem ser equitativos.

Embora reconheçamos que há muitas interpretações diferentes de equidade, consideramos que esta tem uma dimensão substantiva e processual. A dimensão substantiva implica um compromisso com: a garantia de uma distribuição equitativa e justa dos benefícios e dos custos, bem como de inexistência de enviesamentos injustos, discriminação e estigmatização contra pessoas e grupos. Se for possível evitar os enviesamentos, os sistemas de IA podem até aumentar a equidade social. A igualdade de oportunidades em termos de acesso à educação, aos bens e serviços e à tecnologia deve ser igualmente promovida. Além disso, a utilização de sistemas de IA nunca deverá levar a que os utilizadores (finais) sejam iludidos ou prejudicados na sua liberdade de escolha. Além disso, a equidade implica que os profissionais no domínio da IA devem respeitar o princípio da proporcionalidade entre os meios e os fins, e analisar cuidadosamente a forma de equilibrar os interesses e objetivos em causa. A dimensão processual da equidade implica uma possibilidade de contestar e procurar vias de recurso eficazes contra as decisões tomadas por sistemas de IA e pelos seres humanos que os utilizam. Para o efeito, a entidade responsável pela decisão deve ser identificável e os processos decisórios explicáveis.²⁷⁰

Exsurge, nesse ponto, a necessidade de que os algoritmos sejam auditáveis

²⁶⁸ COMISSÃO EUROPEIA. **Orientações éticas para uma IA de Confiança**. Junho de 2018. Disponível em: <http://escola.mpu.mp.br/servicos-academicos/atividades-academicas/inovaescola/curadoria/3-ciclo-de-debates/inteligencia-artificial-e-internet-das-coisas-oportunidades-e-desafios/ethicsguidelinesfortrustworthyai-ptpdf.pdf>. Acesso em: 11 jan. 2022, p. 18.

²⁶⁹ BECK, Ulrich. **Sociedade de risco: rumo a uma outra modernidade**. 2. ed., São Paulo: Editora 34, 2011.

²⁷⁰ COMISSÃO EUROPEIA. **Orientações éticas para uma IA de Confiança**. Junho de 2018. Disponível em: <http://escola.mpu.mp.br/servicos-academicos/atividades-academicas/inovaescola/curadoria/3-ciclo-de-debates/inteligencia-artificial-e-internet-das-coisas-oportunidades-e-desafios/ethicsguidelinesfortrustworthyai-ptpdf.pdf>. Acesso em: 11 jan. 2022, p. 18.

e que a tomada de decisão possa ser compreendida pelos interessados. Simultaneamente, deve-se assegurar que os interesses empresariais do programador não sejam comprometidos, o que, por si só, já cria desafios jurídicos a serem desenvolvidos pela doutrina e pelo legislador. Compreende-se, ainda, que a terminologia “princípio da justiça” não foi a mais adequada, mormente considerando a multiplicidade de conceitos e interpretações que tal vocábulo pode denotar. Analisando a ideia de justiça desenvolvida pelas Diretrizes, pode-se concluir que se quis fazer uma amálgama em referência a princípios de equidade, liberdade, proporcionalidade, solidariedade social e devido processo legal.²⁷¹

Por fim, o documento europeu trata do princípio da explicabilidade, que se refere a prestação de justificativa, que é essencial para construir e manter a confiança dos usuários dos sistemas de IA. Isso significa que os processos precisam ser transparentes, a capacidade e o propósito dos sistemas de IA devem ser comunicados abertamente, e as decisões - na medida possível - devem ser justificadas àqueles afetados direta e indiretamente pelo sistema. Sem essas informações, uma decisão não pode ser devidamente contestada.

Nem sempre é possível explicar por que razão um modelo gerou determinado resultado ou decisão (e que combinação de fatores de entrada contribuiu para esse efeito). Estes casos são designados por algoritmos de «caixa negra» e exigem especial atenção. Nessas circunstâncias, podem ser necessárias outras medidas da explicabilidade (p. ex., a rastreabilidade, a auditabilidade e a comunicação transparente sobre as capacidades do sistema), desde que o sistema, no seu conjunto, respeite os direitos fundamentais. O grau de necessidade da explicabilidade depende em grande medida do contexto e da gravidade das consequências de um resultado errado ou inexato.²⁷²

Não se pode ignorar que a existência da *black box* da IA, compreendida como a ausência de conhecimento acerca da tomada de uma saída ou decisão específica (ou de quais fatores contribuíram para isso), pode desencadear dificuldades. Ainda assim, o respeito aos direitos fundamentais exige que se adote uma postura de explicabilidade, com rastreabilidade e comunicação transparente acerca das capacidades e limitações conhecidas do sistema até então. Trata-se de uma análise contextual, que verifica o funcionamento, os dados, o estado da arte e os objetivos

²⁷¹ SILVA, Gabriela Buarque Pereira; EHRHARDT JÚNIOR, Marcos. Diretrizes éticas para a inteligência artificial confiável na União Europeia e a regulação jurídica no Brasil. *In: Revista IBERC*, v. 3, n. 3, p. 1-28, 2020. Disponível em: <https://revistaiberc.responsabilidadecivil.org/iberc/article/view/133/105>. Acesso em: 11 jan. 2022.

²⁷² COMISSÃO EUROPEIA. **Orientações éticas para uma IA de Confiança**. Junho de 2018. Disponível em: <http://escola.mpu.mp.br/servicos-academicos/atividades-academicas/inovaescola/curadoria/3-ciclo-de-debates/inteligencia-artificial-e-internet-das-coisas-oportunidades-e-desafios/ethicsguidelinesfortrustworthyai-ptpdf.pdf>. Acesso em: 11 jan. 2022, p. 19.

usualmente visados pelo mecanismo tecnológico.²⁷³

Inevitavelmente, tensões podem surgir da interação entre tais princípios, bem como com outros princípios, os direitos fundamentais e os Direitos Humanos, elencados pelo ordenamento jurídico brasileiro, não havendo uma resposta apriorística e fixa para a resolução de tais impasses. Nesse ponto, Alexy²⁷⁴ leciona no sentido de que os princípios são normas que ordenam que algo seja realizado na maior medida possível dentro das possibilidades jurídicas e fáticas existentes. Trata-se de analisar os fatores do contexto e ponderar quais interesses devem prevalecer. Em observância ao desenvolvimento do Estado Democrático de Direito é imprescindível que a percepção e a resolução de tais tensões sejam feitas de modo dialógico com todos os interessados.

3.3.3 Requisitos essenciais para a concretização de uma Inteligência Artificial de confiança

As orientações éticas para uma IA de Confiança, estabelecidas pela Comissão Europeia, fornecem orientações sobre a aplicação e a concretização de uma IA de confiança, através de uma lista de sete requisitos a cumprir, com base nos princípios supramencionados.

A partir das orientações éticas da União Europeia, elencam-se de modo exemplificativo os requisitos que devem ser observados para que o desenvolvimento da inteligência artificial seja confiável: a) ação e supervisão humanas; b) solidez técnica e segurança; c) privacidade e governança de dados; d) transparência; e) diversidade, não discriminação e equidade; f) bem-estar social e ambiental; g) e responsabilização²⁷⁵.

Os princípios descritos no capítulo I devem ser traduzidos em requisitos concretos para alcançar uma IA de confiança. Estes requisitos são aplicáveis a diversas partes interessadas participantes no ciclo de vida dos sistemas de IA: criadores, implantadores e utilizadores finais, bem como a sociedade em

²⁷³ SILVA, Gabriela Buarque Pereira; EHRHARDT JÚNIOR, Marcos. Diretrizes éticas para a inteligência artificial confiável na União Europeia e a regulação jurídica no Brasil. *In: Revista IBERC*, v. 3, n. 3, p. 1-28, 2020. Disponível em: <https://revistaiberc.responsabilidadecivil.org/iberc/article/view/133/105>. Acesso em: 11 jan. 2022.

²⁷⁴ ALEXY, Robert. **Teoria dos direitos fundamentais**. Tradução de Virgílio Afonso da Silva. 2. ed., São Paulo: Malheiros, 2015, p. 588.

²⁷⁵ COMISSÃO EUROPEIA. **Orientações éticas para uma IA de Confiança**. Junho de 2018. Disponível em: <http://escola.mpu.mp.br/servicos-academicos/atividades-academicas/inovaescola/curadoria/3-ciclo-de-debates/inteligencia-artificial-e-internet-das-coisas-oportunidades-e-desafios/ethicsguidelinesfortrustworthyai-ptpdf.pdf>. Acesso em: 11 jan. 2022, p. 20-21.

geral. Designamos por criadores aqueles que investigam, concebem e/ou desenvolvem os sistemas de IA. Por implantadores, entendemos as organizações públicas ou privadas que utilizam os sistemas de IA nos seus processos empresariais e para oferecer produtos e serviços a outros. Os utilizadores finais são as pessoas que interagem de forma direta ou indireta com o sistema de IA. Por último, a sociedade em geral engloba todas as outras pessoas que são direta ou indiretamente afetadas pelos sistemas de IA.²⁷⁶

A União Europeia, ao consolidar os princípios, requisitos e componentes para o justo e correto desenvolvimento de sistemas de IA, fomentou a discussão ético-filosófica que permeia o tema nos demais países e organizações.

Do mesmo modo, a Portaria GM nº 4.617/2021 traz alguns requisitos a serem seguidos:

- Estimular a produção de uma IA ética financiando projetos de pesquisa que visem a aplicar soluções éticas, principalmente nos campos de equidade/não-discriminação (fairness), responsabilidade/prestação de contas (accountability) e transparência (transparency), conhecidas como a matriz FAT.
- Estimular parcerias com corporações que estejam pesquisando soluções comerciais dessas tecnologias de IA ética.
- Estabelecer como requisito técnico em licitações que os proponentes ofereçam soluções compatíveis com a promoção de uma IA ética (por exemplo, estabelecer que soluções de tecnologia de reconhecimento facial adquiridas por órgãos públicos possuam um percentual de falso positivo abaixo de determinado limiar).
- Estabelecer, de maneira multissetorial, espaços para a discussão e definição de princípios éticos a serem observados na pesquisa, no desenvolvimento e no uso da IA.
- Mapear barreiras legais e regulatórias ao desenvolvimento de IA no Brasil e identificar aspectos da legislação brasileira que possam requerer atualização, de modo a promover maior segurança jurídica para o ecossistema digital.
- Estimular ações de transparência e de divulgação responsável quanto ao uso de sistemas de IA, e promover a observância, por tais sistemas, de direitos humanos, de valores democráticos e da diversidade.
- Desenvolver técnicas para identificar e tratar o risco de viés algorítmico.
- Elaborar política de controle de qualidade de dados para o treinamento de sistemas de IA.
- Criar parâmetros sobre a intervenção humana em contextos de IA em que o resultado de uma decisão automatizada implica um alto risco de dano para o indivíduo.
- Incentivar a exploração e o desenvolvimento de mecanismos de revisão apropriados em diferentes contextos de utilização de IA por organizações privadas e por órgãos públicos.
- Criar e implementar melhores práticas ou códigos de conduta com relação à coleta, implantação e uso de dados, incentivando as organizações a melhorar sua rastreabilidade, resguardando os direitos legais.
- Promover abordagens inovadoras para a supervisão regulatória (por exemplo, sandboxes e hubs regulatórios).²⁷⁷

²⁷⁶ COMISSÃO EUROPEIA. **Orientações éticas para uma IA de Confiança**. Junho de 2018. Disponível em: <http://escola.mpu.mp.br/servicos-academicos/atividades-academicas/inovaescola/curadoria/3-ciclo-de-debates/inteligencia-artificial-e-internet-das-coisas-oportunidades-e-desafios/ethicsguidelinesfortrustworthyai-ptpdf.pdf>. Acesso em: 11 jan. 2022, p. 20.

²⁷⁷ BRASIL. **Portaria GM nº 4.617, de 6 de abril de 2021**. Institui a Estratégia Brasileira de Inteligência Artificial e seus eixos temáticos. Disponível em: <https://www.in.gov.br/en/web/dou/-/portaria-gm-n->

A verificação de tais requisitos demanda que haja pesquisa acerca dos sistemas de IA, com divulgação de resultados e abertura de questões ao público. Essa visão estratégica está associada aos valores da IA sólida, robusta e confiável. A partir desses requisitos importantes para o desenvolvimento e uso de soluções e sistemas de IA, busca-se o desenvolvimento de uma IA ética, que segue os desafios postos ao Direito na contemporaneidade.

Passa-se, assim, a analisar cada um dos requisitos propostos pela Comissão Europeia nas orientações éticas para uma IA de Confiança.

a) Ação e supervisão humana:

Quanto à supervisão humana, esta assegura que um sistema de IA não prejudicará a autonomia humana ou causará outros efeitos adversos. A supervisão pode ser alcançada por meio de mecanismos de governança: a abordagem humana no ciclo de vida do sistema, o humano no ciclo de vida do sistema ou o humano no comando do sistema.

A IA deve ser "antropocêntrica e antropogénica (centrada no, e produzida pelo, ser humano)" e realça também a necessidade de as tecnologias "de alto risco" decorrentes da IA -- como a tecnologia de autoaprendizagem, que inclui, por exemplo, a condução autónoma -- serem "concebidas de forma a permitir a supervisão humana a qualquer momento".²⁷⁸

Os sistemas de IA devem apoiar a autonomia e a tomada de decisões dos seres humanos, tal como prescrito pelo princípio de respeito da autonomia humana. Isto exige que os sistemas de IA funcionem como facilitadores de uma sociedade democrática, próspera e equitativa, apoiando a ação do utilizador e a promoção dos direitos fundamentais, e que permitam a supervisão humana.

Cabe salientar a questão do direito de não estar sujeito a uma decisão baseada unicamente no processamento automatizado quando existirem efeitos legais sobre os usuários ou especificidades do caso concreto que demandem uma análise mais acurada. A supervisão humana deve ser garantida, ajudando a eliminar eventuais distorções.²⁷⁹

4.617-de-6-de-abril-de-2021-*313212172. Acesso em: 13 jan. 2022.

²⁷⁸ PARLAMENTO Europeu aprova maior regulamentação de Inteligência Artificial na UE. **In: TSF rádio notícias**, 2020. Disponível em: <https://www.tsf.pt/futuro/parlamento-europeu-aprova-maior-regulamentacao-de-inteligencia-artificial-na-ue-12944972.html>. Acesso em: 13 jan. 2022.

²⁷⁹ SILVA, Gabriela Buarque Pereira; EHRHARDT JÚNIOR, Marcos. Diretrizes éticas para a inteligência artificial confiável na União Europeia e a regulação jurídica no Brasil. **In: Revista IBERC**, v. 3, n. 3, p.

Também pode-se encontrar a previsão do requisito da supervisão humana na Portaria GM nº 4.617/2021, que define a estrutura necessária para a realização da supervisão:

(ii) Estrutura de supervisão (oversight): Estruturas de supervisão normalmente incluem um ou mais representantes legitimados pelo Estado que possuem instrumentos para garantir a aplicabilidade da lei (enforcement), assim como a recomendação de boas práticas e outras salvaguardas. Boas práticas indicam ser desejável que um mecanismo de supervisão bem estruturado inclua elementos como uma autoridade supervisora independente; a obtenção de autorização prévia para a atividade de vigilância (i.e. legalidade); o monitoramento do uso das tecnologias em questão; e a existência de remédios jurídicos eficazes para endereçar eventuais abusos.²⁸⁰

Ademais, as Inteligências Artificiais devem ser concebidas de forma a permitir a supervisão humana a qualquer momento. Contextualizando esta premissa com o cenário brasileiro, pode-se vislumbrar um longo caminho a ser percorrido, sobretudo após a retirada do artigo 20 da Lei Geral de Proteção de Dados Pessoais (lei 13.709/2018), por meio de veto presidencial, da previsão da obrigatoriedade de revisão por pessoa humana das decisões automatizadas.

Todavia há uma problematização no que concerne a supervisão por pessoa humana que advém do Regulamento Geral sobre a Proteção de Dados (RGPD), regulamentação implementada na União Europeia. Segundo o artigo 22 da RGPD:

Art. 22. O titular dos dados tem o direito de **não ficar sujeito a nenhuma decisão tomada exclusivamente com base no tratamento automatizado**, incluindo a **definição de perfis**, que produza efeitos na sua esfera jurídica ou que o afete significativamente de forma similar.
O n. 1 não se aplica se a decisão:
a. For necessária para a celebração ou a execução de um contrato entre o titular dos dados e um responsável pelo tratamento;
b. For autorizada pelo direito da União ou do Estado-Membro a que o responsável pelo tratamento estiver sujeito, e na qual estejam igualmente previstas medidas adequadas para salvaguardar os direitos e liberdades e os legítimos interesses do titular dos dados; ou
c. For baseada no consentimento explícito do titular dos dados.²⁸¹

Sendo assim, a regra elencada pelo artigo se vincula a possibilidade de rejeição da utilização de sistemas inteligentes, não estando, em regra, os titulares dos dados em questão, reféns do sistema. Ainda desta decisão automatizada, o próprio

1-28, 2020. Disponível em: <https://revistaiberc.responsabilidadecivil.org/iberc/article/view/133/105>. Acesso em: 13 jan. 2022.

²⁸⁰ BRASIL. **Portaria GM nº 4.617, de 6 de abril de 2021**. Institui a Estratégia Brasileira de Inteligência Artificial e seus eixos temáticos. Disponível em: https://www.in.gov.br/en/web/dou/-/portaria-gm-n-4.617-de-6-de-abril-de-2021-*313212172. Acesso em: 13 jan. 2022.

²⁸¹ UNIÃO EUROPEIA. **Regulamento (UE) 2016/679 do Parlamento Europeu e do Conselho de 27 de abril de 2016**. Regulamento Geral sobre a Proteção de Dados. Disponível em: <https://eur-lex.europa.eu/legal-content/PT/TXT/?uri=celex%3A32016R0679>. Acesso em: 13 jan. 2022.

regulamento garante uma contestação analisada por pessoa humana.

Desta forma, tomando como exemplo o fato de que as instituições financeiras utilizam de uma ferramenta de IA para criação de perfis de indivíduos na análise de concessão de empréstimo, e assim, de acordo com seu resultado, autorizam ou não a concessão do empréstimo, pode o titular dos dados, negar-se a tê-los analisados por uma ferramenta deste tipo.

Diferentemente, no caso específico da legislação brasileira, a formatação do artigo 20, da LGPD estipula que:

Art. 20. O titular dos dados tem direito a solicitar a revisão de decisões tomadas unicamente com base em tratamento automatizado de dados pessoais que afetem seus interesses, incluídas as decisões destinadas a definir o seu perfil pessoal, profissional, de consumo e de crédito ou os aspectos de sua personalidade.²⁸²

Cabe frisar, ainda, que isso inclui o poder de decisão de não usar o sistema de IA em uma situação particular, de estabelecer níveis de descrição humana durante o uso do sistema, e/ou de garantir a capacidade de anular uma decisão tomada por um sistema.

Dessa maneira, a tomada de decisão judicial com o auxílio de IA pode ser mais célere e eficiente sem que necessariamente haja perda da qualidade deliberativa, uma vez que a habilidade de chegar a uma decisão e implementá-la nos algoritmos é importante na tomada de decisões através de IA, desde que garantido o posterior controle e a supervisão por humanos (juizador e sua equipe técnica), eventuais vieses e ou saídas incorretas (sem acurácia) podem ser tratados antes de ser efetivamente adotada a decisão sugerida pela IA judicial, devendo o erro da IA ser registrado no sistema (art. 25, Parágrafo único, V, da Res. 332/2020) para que exista a transparência da decisão tomada.²⁸³

b) Solidez técnica e segurança

Conforme disposto nas diretrizes éticas para IA de Confiança, um componente crucial para que a IA de confiança se torne realidade é a solidez técnica, que está estreitamente ligada ao princípio da prevenção de danos. A solidez técnica exige que os sistemas de IA sejam desenvolvidos seguindo uma abordagem de prevenção dos riscos e de forma a que se comportem fielmente conforme o previsto, minimizando os

²⁸² BRASIL. **Lei nº 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: 13 jan. 2022.

²⁸³ PEIXOTO, Fabiano Hartmann; SILVA, Roberta Zumblick Martins da. **Inteligência Artificial e Direito**. Vol. 1. Curitiba: Alteridade Editora, 2019, p. 40.

danos não intencionais e inesperados, e prevenindo os danos inaceitáveis. Tal deverá também ser aplicado a eventuais alterações do ambiente em que operam ou à presença de outros agentes (humanos e artificiais) que possam interagir com o sistema de forma antagônica. Além disso, deve assegurar-se a integridade física e mental dos seres humanos.²⁸⁴

Em outras palavras, os sistemas devem ser protegidos contra a vulnerabilidade de exploração por outras máquinas, como por exemplo, hackers. Estes ataques podem afetar diretamente dados submetidos ao sistema, podendo ocasionar a adulteração deste ou mesmo o vazamento. Aliás, não estarão em perigo somente os dados como também o próprio sistema programado. A infiltração de sistemas não desejados pode acarretar no corrompimento da própria inteligência artificial, afetando todo seu processo decisório.

O nível de medidas de segurança requeridas em determinado sistema de IA deverá ser proporcional ao nível da magnitude do risco apresentado pela máquina. O risco será maior na medida em que a máquina tiver menor precisão em seu funcionamento, isto é, menor capacidade de fazer julgamentos e classificações corretas com base nos dados ou modelos.²⁸⁵

As Diretrizes estipulam, ainda, que quando previsões imprecisas ocasionais não puderem ser evitadas, é importante que o sistema possa indicar a probabilidade de ocorrência desses erros, o que concretiza a noção de boa-fé objetiva. É de extrema relevância que haja um mínimo controle sobre aquilo que se produz, para que seja possível, de modo transparente, informar os interessados sobre as limitações verificadas no produto ou serviço.²⁸⁶

c) Privacidade e governação dos dados

Estreitamente ligado ao princípio de prevenção de danos está o direito à privacidade, um direito fundamental que é particularmente afetado pelos sistemas de

²⁸⁴ COMISSÃO EUROPEIA. **Orientações éticas para uma IA de Confiança**. Junho de 2018. Disponível em: <http://escola.mpu.mp.br/servicos-academicos/atividades-academicas/inovaescola/curadoria/3-ciclo-de-debates/inteligencia-artificial-e-internet-das-coisas-oportunidades-e-desafios/ethicsguidelinesfortrustworthyai-ptpdf.pdf>. Acesso em: 13 jan. 2022, p. 23.

²⁸⁵ SILVA, Gabriela Buarque Pereira; EHRHARDT JÚNIOR, Marcos. Diretrizes éticas para a inteligência artificial confiável na União Europeia e a regulação jurídica no Brasil. **In: Revista IBERC**, v. 3, n. 3, p. 1-28, 2020. Disponível em: <https://revistaiberc.responsabilidadecivil.org/iberc/article/view/133/105>. Acesso em: 13 jan. 2022.

²⁸⁶ SILVA, Gabriela Buarque Pereira; EHRHARDT JÚNIOR, Marcos. Diretrizes éticas para a inteligência artificial confiável na União Europeia e a regulação jurídica no Brasil. **In: Revista IBERC**, v. 3, n. 3, p. 1-28, 2020. Disponível em: <https://revistaiberc.responsabilidadecivil.org/iberc/article/view/133/105>. Acesso em: 13 jan. 2022.

IA. A prevenção da ameaça à privacidade também exige uma governação adequada dos dados, que assegure a qualidade e a integridade dos dados utilizados, a sua relevância para o domínio em que os sistemas de IA serão implantados, os seus protocolos de acesso e a capacidade de tratar os dados de modo a proteger a privacidade.²⁸⁷

Os sistemas de IA devem garantir a privacidade e a proteção de dados durante todo o ciclo da vida, desde o momento em que as informações são fornecidas, bem como aquelas produzidas pela máquina. Objetivando a confiabilidade, deve ser assegurado que ao recolhimento de dados, estes não sejam utilizados de maneira discriminatória ou de forma ilegal e injusta. No ordenamento brasileiro, trata-se de assegurar a observância do art. 5º, X, da Constituição Federal.

Salientando que além de ser um Direito Fundamental, a garantia da privacidade e proteção dos dados está diretamente ligada aos Direitos Humanos. Neste mesmo sentido, Ingo Wolfgang Sarlet²⁸⁸, traz em seu livro “Eficácia dos Direitos Fundamentais”, leciona que “o direito fundamental à privacidade é faceta da própria dignidade da pessoa humana, sob valor constitucional, em seu aspecto subjetivo, conforme Declaração Universal dos Direitos Humanos”.

É indiscutível a obrigação do respeito aos direitos de proteção de dados e privacidade, abordados na própria Carta Magna, nas implantações de Inteligências Artificiais no Poder Público.

Além disso, a qualidade do fornecimento é requisito fundamental para o desempenho do sistema, não havendo possibilidade de haver erros, inexatidões ou enganos no momento da entrega, caso contrário, em casos de sistemas não supervisionados, estas informações podem alterar o comportamento da máquina em relação ao procedimento adotado e, conseqüentemente, sua decisão.

Em relação à governança dos dados, em qualquer organização que trabalhe com dados devem ser adotados protocolos para acesso, devendo indicar quem terá permissão para ingresso no sistema e em que circunstâncias esse acesso será permitido, garantindo uma maior segurança e controle das informações fornecidas.

²⁸⁷ COMISSÃO EUROPEIA. **Orientações éticas para uma IA de Confiança**. Junho de 2018. Disponível em: <http://escola.mpu.mp.br/servicos-academicos/atividades-academicas/inovaescola/curadoria/3-ciclo-de-debates/inteligencia-artificial-e-internet-das-coisas-oportunidades-e-desafios/ethicsguidelinesfortrustworthyai-ptpdf.pdf>. Acesso em: 13 jan. 2022, p. 24.

²⁸⁸ SARLET, Ingo. **Eficácia dos Direitos Fundamentais**. Porto Alegre: Ed. Livraria do Advogado, 2010.

Nesse sentido, impõe-se o respeito pela privacidade, qualidade, integridade e acesso aos dados. Os autores Erick Lucena e Marcos Ehrhardt Júnior, apresentam algumas estratégias que podem ser traçadas para a tutela da privacidade:

[...] a primeira delas seria o “direito de oposição”, que, de forma individual ou coletiva, funcionaria como uma negativa à coleta e circulação de informações pessoais em determinadas formas. b) O “direito de não saber” é a segunda estratégia de tutela da privacidade, podendo ser tratado como decorrente do primeiro. Surgido em relação a dados de saúde, passa a ser estendido contra as formas de marketing direto que invadem a esfera privada do indivíduo com informações não solicitadas e não desejadas. c) Outra estratégia é tornar mais clara a finalidade da coleta de dados. A legitimidade aqui é condicionada à comunicação preventiva ao interessado sobre o motivo da coleta e o destino dos dados coletados. d) Por último, o “direito ao esquecimento”, “prevendo-se que algumas categorias de informações devam ser destruídas, ou conservadas somente em forma agregada e anônima, uma vez que tenha sido atingida a finalidade para a qual foram coletadas” ou ainda, “depois de transcorrido um determinado lapso de tempo.”²⁸⁹

A qualidade e a integridade dos dados são cruciais ao desempenho dos sistemas de IA, considerando que quando os dados coletados contêm vieses, imprecisões ou falhas estruturais, tais erros serão reproduzidos pela tecnologia. No mesmo sentido, a alimentação de dados maliciosos pode mudar o comportamento do sistema, especialmente no que tange à autoaprendizagem.

d) Transparência

Este requisito está estreitamente relacionado com o princípio da explicabilidade e abrange a transparência dos elementos relevantes para um sistema de IA: os dados, o sistema e os modelos de negócio.

Para ser possível uma relação de confiança entre a máquina e o humano, é basilar o exercício da transparência entre as partes, sendo assim, um dos grandes desafios enfrentados pela IA são os chamados *blackbox* já tratados no primeiro capítulo. Os sistemas têm o dever de documentar todas as ações tomadas para que assim permita a rastreabilidade e a lucidez dos programas, facilitando a identificação dos motivos pelos quais um processo decisório foi ineficaz ou equivocado, ajudando na manutenção para que não ocorram novos erros.

No que se refere a rastreabilidade, os conjuntos de dados e os processos que produzem a decisão do sistema de IA incluindo os processos de recolha e etiquetagem dos dados, bem como os algoritmos utilizados, devem ser documentados da melhor forma possível para permitir a rastreabilidade e um aumento da transparência. Isto

²⁸⁹ PEIXOTO, Erick Lucena Campos; EHRHARDT JÚNIOR, Marcos. Breves notas sobre a resignificação da privacidade. *In: Revista Brasileira de Direito Civil – RBDCivil*, Belo Horizonte, v. 16, p. 35-56, abr./jun. 2018, p. 44.

também se aplica às decisões tomadas pelo sistema de IA. Deste modo, é possível identificar os motivos por que uma decisão de IA foi errada, o que, por sua vez, poderá ajudar a evitar erros futuros. A rastreabilidade facilita, assim, a auditabilidade e a explicabilidade.²⁹⁰

Não obstante, a transparência auxilia na compreensão e consequente rastreamento pelo ser humano nas tomadas de decisão da Inteligência Artificial, sendo auxiliar também na revisão pelo ser humano para identificar a origem do erro e assim evitar que se replique.

Além disso, visando a devida comunicação, é necessário haver uma informação aos usuários de que irão interagir com uma Inteligência Artificial e estão facultados a solicitar intervenção humana, garantindo a aplicação dos direitos fundamentais da publicidade e informação.

Conforme dispõe a Portaria GM nº 4.617/2021,

A transparência constitui elemento importante de estruturas de governança de IA, seja no que se refere à informação quanto à interação com sistemas de IA (disclosure), seja no que tange à ideia de explicabilidade de decisões tomadas por sistemas autônomos, conforme debatido anteriormente. Do ponto de vista procedimental, a ideia de transparência pode ser traduzida pela adoção de metodologias transparentes e auditáveis quanto ao desenvolvimento dos sistemas de IA, quanto às fontes de dados e quanto aos procedimentos e documentação do projeto em questão.²⁹¹

Como identificado, as Inteligências Artificiais utilizadas pelo poder público funcionam a partir de precedente, ou seja, seus algoritmos são treinados com dados anteriormente julgados que podem, de certa forma, serem tendenciosos, podendo, eventualmente, gerar segregações sociais em uma escala cada vez mais automatizada, dificultando o ingresso das máquinas ao poder judiciário ou até mesmo a qualquer empresa.

A transparência é um dos requisitos de grande importância para a construção de uma IA de confiança, tanto é que o Regulamento Geral de Proteção de Dados traz uma seção inteira tratando da transparência, denominado “Transparência e regras para o exercício dos direitos dos titulares dos dados”.²⁹²

²⁹⁰ COMISSÃO EUROPEIA. **Orientações éticas para uma IA de Confiança**. Junho de 2018. Disponível em: <http://escola.mpu.mp.br/servicos-academicos/atividades-academicas/inovaescola/curadoria/3-ciclo-de-debates/inteligencia-artificial-e-internet-das-coisas-oportunidades-e-desafios/ethicsguidelinesfortrustworthyai-ptpdf.pdf>. Acesso em: 13 jan. 2022, p. 24-25.

²⁹¹ BRASIL. **Portaria GM nº 4.617, de 6 de abril de 2021**. Institui a Estratégia Brasileira de Inteligência Artificial e seus eixos temáticos. Disponível em: <https://www.in.gov.br/en/web/dou/-/portaria-gm-n-4.617-de-6-de-abril-de-2021-313212172>. Acesso em: 13 jan. 2022.

²⁹² UNIÃO EUROPEIA. **Regulamento (UE) 2016/679 do Parlamento Europeu e do Conselho de 27 de abril de 2016**. Regulamento Geral sobre a Proteção de Dados. Disponível em: <https://eur->

Desta forma, a opacidade na tomada de decisões automatizadas proporciona a população uma insegurança quanto ao uso destas, visto que não há possibilidade de identificação quanto aos meios e procedimentos adotados para obtenção de um resultado.

e) Diversidade, não discriminação e equidade

A inclusão e a diversidade têm de estar presentes em todo o ciclo de vida do sistema de IA para que a IA de confiança se torne uma realidade. Além da consideração e do envolvimento de todas as partes interessadas ao longo do processo, tal implica também que a igualdade de acesso mediante processos inclusivos e a igualdade de tratamento sejam asseguradas. Este requisito está estreitamente relacionado com o princípio da equidade.²⁹³

Quando se fala na aplicação de precedentes em máquinas de inteligência artificial, temos que considerar que eventuais erros, parcialidades e equívocos cometidos na esfera humana, serão replicados na tomada de decisão autônoma.

Segundo o documento das Orientações Éticas para uma IA de Confiança:

Os conjuntos de dados utilizados pelos sistemas de IA (tanto para treino como para funcionamento) podem ser afetados pela inclusão de desvios históricos inadvertidos, bem como por lacunas e por maus modelos de governação. A manutenção de tais desvios pode dar origem a discriminação e preconceitos (in)diretos não intencionais contra determinados grupos ou pessoas, agravando o preconceito e a marginalização. A exploração intencional de preconceitos já existentes (entre os consumidores) e as práticas de concorrência desleal, tais como a homogeneização dos preços através de conluíus ou da falta de transparência do mercado, também podem causar danos. O enviesamento identificável e discriminatório deve ser eliminado na fase de recolha de dados, sempre que possível. A forma como os sistemas de IA são desenvolvidos (p. ex., a programação de algoritmos) também pode ser afetada por um enviesamento injusto. Tal pode ser combatido mediante a adoção de processos de supervisão para analisar e abordar a finalidade, os condicionalismos, os requisitos e as decisões do sistema de forma clara e transparente. Além disso, o recrutamento de pessoal de diferentes origens, culturas e disciplinas pode assegurar a diversidade de opiniões e deve ser incentivado.²⁹⁴

Ademais, quando não for possível a intervenção nas fases embrionárias do procedimento como no momento do recolhimento de dados, deverá então ser

lex.europa.eu/legal-content/PT/TXT/?uri=celex%3A32016R0679. Acesso em: 13 jan. 2022.

²⁹³ COMISSÃO EUROPEIA. **Orientações éticas para uma IA de Confiança**. Junho de 2018. Disponível em: <http://escola.mpu.mp.br/servicos-academicos/atividades-academicas/inovaescola/curadoria/3-ciclo-de-debates/inteligencia-artificial-e-internet-das-coisas-oportunidades-e-desafios/ethicsguidelinesfortrustworthyai-ptpdf.pdf>. Acesso em: 13 jan. 2022, p. 25.

²⁹⁴ COMISSÃO EUROPEIA. **Orientações éticas para uma IA de Confiança**. Junho de 2018. Disponível em: <http://escola.mpu.mp.br/servicos-academicos/atividades-academicas/inovaescola/curadoria/3-ciclo-de-debates/inteligencia-artificial-e-internet-das-coisas-oportunidades-e-desafios/ethicsguidelinesfortrustworthyai-ptpdf.pdf>. Acesso em: 13 jan. 2022, p. 25-26.

combatido no procedimento da supervisão humana, analisando a falha e apresentando relatório para que seja corrigido, preservando pareceres futuros. Nesta perspectiva, conforme já foi demonstrado em capítulo anterior, nota-se que esta discrepância de valores já foi identificada no sistema adotado pelos Estados Unidos em casos criminais, o chamado Sistema COMPAS.

Como mencionado anteriormente, o método adotado determina, através de uma pontuação obtida por algoritmos, o possível índice de reincidência e periculosidade de determinado sujeito. A polêmica reside nas decisões claramente tendenciosas e na impossibilidade de questionamento do resultado, uma vez que é propriedade de uma empresa privada e se trata de um segredo comercial.

Os vieses discriminatórios devem ser tolhidos já na fase de coleta, de modo que os critérios a serem utilizados no processamento da IA já estejam livres de tais falhas. É importante, assim, que a base de dados seja inclusiva no que tange a diversas culturas e origens. Tais problemas também podem ser mitigados com supervisões que analisem finalidade, restrições, requisitos e decisões do sistema de maneira coerente e transparente.

É fácil perceber que, se forem utilizados no modelo estatístico dados com alto potencial discriminatório, tais como dados raciais, étnicos ou de orientação sexual, haverá um grande risco de que a decisão que resultará do processo automatizado (output) também seja discriminatória. Esses dados são os chamados dados sensíveis, cujo processamento é limitado pelas legislações de proteção de dados de vários países, assim como pelo Regulamento Europeu de Dados Pessoais. Em segundo lugar, é preciso observar que o próprio método utilizado nas decisões automatizadas – por meio da classificação e seleção dos indivíduos – gera um risco de se produzirem resultados discriminatórios, ainda que de forma não intencional. Isto pode ocorrer porque, na discriminação estatística, teoria econômica que se tornou conhecida a partir dos textos de Edmund Phelps (1972) e Kenneth Arrow (1973), os indivíduos são diferenciados com base em características prováveis de um grupo, no qual esse indivíduo é classificado. Essa prática se baseia em métodos estatísticos, que associam esses atributos a outras características, cuja identificação pelo tomador de decisão é mais difícil, como nível de renda, risco de inadimplência, produtividade no trabalho, etc. Nesse contexto, é possível a ocorrência da discriminação por erro estatístico, o que decorreria tanto de dados incorretamente capturados como também de modelo estatístico de bases científicas frágeis.²⁹⁵

Projetar modelos é uma tarefa permeada por subjetividade, ainda que seu grau possa variar de acordo com o tipo de processo a ser modelado. Em todo caso, para que um modelo tenha potencial de causar grande prejuízo para um grupo de

²⁹⁵ DONEDA, Danilo Cesar Maganhoto; MENDES, Laura Schertel; SOUZA, Carlos Affonso Pereira de; ANDRADE, Norberto Nuno Gomes de. Considerações iniciais sobre inteligência artificial, ética e autonomia pessoal. *In: Pensar - Revista de Ciências Jurídicas*, Fortaleza, vol. 23, n. 4, p. 1-17, out./dez. 2018, p. 5.

peessoas, é preciso uma conjunção de fatores. Para Cathy O'Neil²⁹⁶, um modelo se torna uma “arma de destruição matemática” quando ele é (i) opaco, (ii) possui a capacidade de crescer exponencialmente e (iii) é projetado para operar prejudicar aqueles a ele sujeitos.

Especialmente quando utilizados pela Administração Pública, modelos têm um grande potencial de se tornarem danosos. Afinal, por pressuposto, seus usos serão estendidos a um grande número de pessoas e, junto disso, modelos apresentam uma “tendência à inescrutabilidade”. Dessa forma, os dois primeiros requisitos estarão, em boa parte dos casos, automaticamente preenchidos ou em sua iminência. O terceiro requisito se torna, portanto, fundamental, uma vez que a forma como os modelos são projetados e os propósitos para os quais eles serão utilizados impactarão diretamente no seu potencial lesivo.²⁹⁷

Com relação à sua opacidade, em um Estado democrático e de direito, ela será sempre um desafio. O simples fato de um modelo ser escrito através de uma notação matemática e/ou por meio de uma linguagem computacional já o torna inacessível à maior parte das pessoas. Esse problema se torna ainda mais pertinente, quando se trata de modelos que utilizam aprendizado de máquina, cujos detalhes de seu modo de funcionamento não são conhecidos nem mesmo por seus programadores.²⁹⁸

Em resumo, dessa forma, diminui-se a chance de que tais algoritmos se tornem máquinas projetadas para arruinar vida de pessoas ou que entrem em ciclos de retroalimentação viciosos. Deve-se observar que um mesmo algoritmo, treinado com os mesmos dados, pode ser mais ou menos danoso, a depender do uso que se faz dele.

f) Bem-estar social e ambiental

Em conformidade com os princípios da equidade e da prevenção de danos, a sociedade em geral, outros seres sensíveis e o ambiente também devem ser considerados partes interessadas ao longo do ciclo de vida da IA. A sustentabilidade e a responsabilidade ecológica dos sistemas de IA devem ser incentivadas e deve ser

²⁹⁶ O'NEIL, Cathy. **Weapons of math destruction**. New York: Broadway Books, 2016.

²⁹⁷ BOEING, Daniel Henrique Arruda. **Ensinando um robô a julgar**: pragmática, discricionariedade e vieses no uso de aprendizado de máquina no judiciário. 2019. 84 fls. Monografia (Graduação em Direito) – Universidade Federal de Santa Catarina, Florianópolis, 2019.

²⁹⁸ BOEING, Daniel Henrique Arruda. **Ensinando um robô a julgar**: pragmática, discricionariedade e vieses no uso de aprendizado de máquina no judiciário. 2019. 84 fls. Monografia (Graduação em Direito) – Universidade Federal de Santa Catarina, Florianópolis, 2019.

promovida a investigação em soluções de IA direcionadas para áreas de interesse global, como, por exemplo, os Objetivos de Desenvolvimento Sustentável. Idealmente, a IA deve ser utilizada em benefício de todos os seres humanos, incluindo as gerações futuras.²⁹⁹

Nesse ponto, o art. 225 da Constituição Federal de 1988 estipula que todos têm direito ao meio ambiente ecologicamente equilibrado, bem de uso comum do povo e essencial à sadia qualidade de vida, impondo-se ao Poder Público e à coletividade o dever de defendê-lo e preservá-lo para as presentes e futuras gerações. Trata-se, portanto, de direito fundamental baseado na noção de solidariedade social.

Também se recomenda que os efeitos sociais oriundos da IA sejam devidamente analisados e monitorados, considerando que a tecnologia se torna cada vez mais presente e invasiva no cotidiano dos indivíduos, às vezes de modo bem sutil. Tal fenômeno pode ensejar modificações nas relações sociais, econômicas e culturais, contribuindo e simultaneamente deteriorando habilidades e costumes sociais. O uso de sistemas de IA deve ser cautelosamente analisado, máxime tendo em vista que até mesmo em contextos eleitorais a tecnologia vem sendo utilizada para influenciar posturas e perspectivas sociais.³⁰⁰

Nesse aspecto, as Orientações Éticas da União Europeia prescrevem:

Além de se avaliar o impacto do desenvolvimento, da implantação e da utilização de um sistema de IA nos indivíduos, também se deverá avaliar esse impacto numa perspectiva societal, tendo em conta o seu efeito nas instituições, na democracia e na sociedade em geral. A utilização de sistemas de IA deve ser cuidadosamente ponderada, em especial em situações relacionadas com o processo democrático, incluindo não só o processo de tomada de decisões políticas, mas também os contextos eleitorais.³⁰¹

Importante também mencionar o que a Portaria GM nº 4.617/2021 dispõe:

As preocupações com a dignidade humana e com a valorização do bem-estar humano devem estar presentes desde a concepção (ethics by design) dessas ferramentas até a verificação de seus efeitos na realidade dos cidadãos. Frise-se que o desenvolvimento de uma Sociedade do Futuro centrada no ser humano é uma das diretrizes adotadas pelo "G20 - Declaração Ministerial

²⁹⁹ COMISSÃO EUROPEIA. **Orientações éticas para uma IA de Confiança**. Junho de 2018. Disponível em: <http://escola.mpu.mp.br/servicos-academicos/atividades-academicas/inovaescola/curadoria/3-ciclo-de-debates/inteligencia-artificial-e-internet-das-coisas-oportunidades-e-desafios/ethicsguidelinesfortrustworthyai-ptpdf.pdf>. Acesso em: 13 jan. 2022, p. 26.

³⁰⁰ SILVA, Gabriela Buarque Pereira; EHRHARDT JÚNIOR, Marcos. Diretrizes éticas para a inteligência artificial confiável na União Europeia e a regulação jurídica no Brasil. *In: Revista IBERC*, v. 3, n. 3, p. 1-28, 2020. Disponível em: <https://revistaiberc.responsabilidadecivil.org/iberc/article/view/133/105>. Acesso em: 13 jan. 2022.

³⁰¹ COMISSÃO EUROPEIA. **Orientações éticas para uma IA de Confiança**. Junho de 2018. Disponível em: <http://escola.mpu.mp.br/servicos-academicos/atividades-academicas/inovaescola/curadoria/3-ciclo-de-debates/inteligencia-artificial-e-internet-das-coisas-oportunidades-e-desafios/ethicsguidelinesfortrustworthyai-ptpdf.pdf>. Acesso em: 13 jan. 2022, p. 27.

sobre Comércio e Economia Digital - Princípios para IA Centrada nos Humanos (2019)⁴⁰" ao tratar de economia digital, de IA e de meios para que as políticas digitais maximizem benefícios e minimizem desafios.³⁰²

Sendo assim, os desenvolvedores da tecnologia devem lidar com as preocupações ambientais, sem negligenciar os impactos que podem advir de sua atuação. O processo de desenvolvimento, bem como toda a sua cadeia de produção e suprimentos, deve ser avaliado por meio da análise de seus recursos e consumo de energia, optando sempre por opções menos prejudiciais. Trata-se da consagração da função social da empresa, que tem o dever de observar as normas cogentes que versam sobre a preservação do meio ambiente.

g) Responsabilização

Por fim, o requisito da responsabilização exige a criação de mecanismos que possibilitem a responsabilização dos sistemas de inteligência artificial por seus resultados.

O requisito de responsabilização complementa os requisitos acima enunciados, estando estreitamente relacionado com o princípio da equidade. Exige que sejam criados mecanismos para garantir a responsabilidade e a responsabilização pelos sistemas de IA e os seus resultados, tanto antes como depois da sua adoção.³⁰³

Essa responsabilização se refere à possibilidade de auditoria, minimização de impactos negativos e reparação. A auditoria se refere à avaliação de algoritmos, dados e processos de design. Muito se argumenta acerca da possível violação de modelos de negócios e propriedade intelectual dos programadores por auditores externos e internos, o que demanda esforços no sentido de compatibilizar o sigilo de tais empresários com a necessidade de auditoria independente. Em sistemas que afetam direitos fundamentais, incluindo aplicações críticas de segurança, é imprescindível que os sistemas possam ser auditados.

Quanto a ética de algoritmos, deve imperar a responsabilidade algorítmica, de modo que, se faz necessário observar os algoritmos como objeto de criação e levar em consideração a intenção, inclusive de qualquer grupo, instituição, agência ou

³⁰² BRASIL. **Portaria GM nº 4.617, de 6 de abril de 2021**. Institui a Estratégia Brasileira de Inteligência Artificial e seus eixos temáticos. Disponível em: https://www.in.gov.br/en/web/dou/-/portaria-gm-n-4.617-de-6-de-abril-de-2021-*313212172. Acesso em: 13 jan. 2022.

³⁰³ COMISSÃO EUROPEIA. **Orientações éticas para uma IA de Confiança**. Junho de 2018. Disponível em: <http://escola.mpu.mp.br/servicos-academicos/atividades-academicas/inovaescola/curadoria/3-ciclo-de-debates/inteligencia-artificial-e-internet-das-coisas-oportunidades-e-desafios/ethicsguidelinesfortrustworthyai-ptpdf.pdf>. Acesso em: 13 jan. 2022, p. 27.

recursos humanos, ou seja, qualquer ator, que possam ter influenciado nos processos de design algorítmico.³⁰⁴

A necessidade de reparação se avulta sob a noção da teoria do risco, com fulcro no art. 927, parágrafo único, do Código Civil, que determina que há obrigação de reparar o dano independentemente de culpa, nos casos especificados em lei, ou quando a atividade normalmente desenvolvida pelo autor do dano implicar, por sua natureza, risco para os direitos de outrem. Trata-se da responsabilidade objetiva, que dispensa verificação de culpa para incidência e se lastreia na necessidade de assegurar à vítima a reparação de seu prejuízo.

Nesse sentido, argumenta Maria Celina Bodin de Moraes³⁰⁵:

Com o passar do tempo, porém, o dever de solidariedade social, o fundamento constitucional da responsabilidade objetiva, sobressairá e aceitar-se-á que seu alcance é amplo o suficiente para abranger a reparação de todos os danos injustamente sofridos, em havendo nexo de causalidade com a atividade desenvolvida, seja ela perigosa ou não. Não se sustentará mais qualquer resquício de culpa, de sanção ou de descumprimento de deveres no fundamento da responsabilidade objetiva. Com efeito, todas são atividades que geram 'risco para os direitos de outrem', como prevê o dispositivo legal.

Deve ser assegurada a capacidade de comunicar as ações ou decisões que contribuem para um determinado resultado do sistema, bem como de responder às consequências desse resultado. A identificação, a avaliação, a comunicação e a minimização dos potenciais impactos negativos dos sistemas de IA são particularmente cruciais para as pessoas (in)diretamente afetadas. Deve disponibilizar-se a devida proteção aos denunciante, às ONG, aos sindicatos ou a outras entidades que denunciem preocupações legítimas com um sistema baseado na IA. O recurso a avaliações de impacto tanto antes como durante o desenvolvimento, a implantação e a utilização de sistemas de IA pode ser útil para minimizar os impactos negativos. Estas avaliações devem ser proporcionadas em relação ao risco colocado pelos sistemas de IA.

Uma preocupação crítica no domínio da IA surge quando uma das componentes da IA de confiança é violada. Muitas das preocupações podem ser

³⁰⁴ MAGRANI, Eduardo. A governança da internet das coisas e a ética da inteligência artificial. *In: Revista Direitos Culturais*, Santo Ângelo, vol. 13, n. 31, p. 153-190, set./dez. 2018. Disponível em: <http://eduardomagrani.com/wp-content/uploads/2019/02/ARTIGO-ETICA-E-AI-PUBLICADO-REVISTA-DE-DIRITEOS-CULTURALS-1.pdf>. Acesso em: 13 jan. 2022.

³⁰⁵ MORAES, Maria Celina Bodin de. A constitucionalização do direito civil e seus efeitos sobre a responsabilidade civil. *In: Direito, Estado e Sociedade*, Rio de Janeiro, n. 29, p. 233-258, 2006, p. 251.

abrangidas pelos princípios e requisitos apresentados nas Orientações Éticas da Comissão da União Europeia, contudo, se sabe que poderão surgir preocupações ainda desconhecidas com a implementação e desenvolvimento da IA. Por este motivo, é necessário construir uma regulamentação sólida, com base em parâmetros éticos, centrado nos Direitos Fundamentais e os Direitos Humanos, bem como reconhecer seu caráter transnacional, o que demonstra a necessidade de se pensar em uma regulamentação que considere a realidade Global da implementação desta tecnologia.

Contudo, cabe salientar que existe uma clara lacuna sobre a responsabilização civil para as questões que envolvem inteligência artificial, por envolver situações mais complexas e imprevisíveis, que transcendem a falha de algum componente do produto, de defeito de fabricação, de negligência do fabricante ou mesmo da insuficiência de testes. Leonardo Netto Parentoni, Rômulo Soares Valentini e Tárik César Oliveira e Alves trazem um ótimo exemplo de uma situação nesse sentido:

Um carro autônomo passou com sucesso por mais de dez mil horas de testes, tanto em ambiente controlado quanto nas ruas de determinada cidade. O suficiente para ser aprovado segundo os mais rigorosos padrões de segurança à época aplicáveis. Foi então colocado em circulação. Até que, determinado dia, esse carro autônomo trafegava normalmente quando, de repente, uma sacola de plástico de supermercado foi arrancada pelo vento das mãos de um comprador e passou voando na frente do carro, que a confundiu com uma criança e, ao desviar, causou um acidente. É lícito e razoável exigir que os testes houvessem antevisto situações tão específicas como essa? Justifica-se transferir esse tipo de risco ao desenvolvedor do produto? Justifica-se prolongar o período de testes quase indefinidamente, para simular até mesmo as situações mais improváveis e, por consequência, atrasar a entrada no mercado de produtos que podem ser muito úteis, socialmente desejáveis e ansiosamente aguardados?³⁰⁶

Tem-se, aqui, aquilo que Andreas Matthias denominou “lacuna de responsabilidade”:

Tradicionalmente, tomamos o operador ou o fabricante da máquina como responsável pelas consequências de sua operação, ou ‘ninguém’ (em casos em que nenhuma falha pessoal possa ser identificada). Agora, pode ser demonstrado que há uma crescente classe de ações de máquinas, onde as formas tradicionais de atribuição de responsabilidade não são compatíveis com nosso senso de justiça e com a estrutura moral da sociedade, porque ninguém tem controle suficiente sobre as ações da máquina para poder assumir a responsabilidade por elas.³⁰⁷

³⁰⁶ PARENTONI, Leonardo Netto; VALENTINI, Rômulo Soares; ALVES, Tárik César Oliveira e. Panorama da regulação da inteligência artificial no Brasil: com ênfase no PLS n. 5.051/2019. *In: Revista Eletrônica do Curso de Direito da UFSM*, v. 15, n. 2 / 2020. Disponível em: <https://periodicos.ufsm.br/revistadireito/article/download/43730/pdf/246146>. Acesso em: 09 fev. 2022.

³⁰⁷ MATTHIAS, Andreas. The responsibility gap: ascribing responsibility for the actions of learning automata. *In: Ethics and Information Technology*, v. 6, n. 3, p. 175-183, 2004.

No entanto, antes de se definir que tipo de responsabilidade se deve utilizar nos casos de danos causados pela inteligência artificial, necessário se faz definir quem são os agentes envolvidos? Nesse sentido, a resolução do Parlamento Europeu, de 20 de outubro de 2020, que contém recomendações à Comissão sobre o regime de Responsabilidade Civil aplicável à Inteligência Artificial (2020/2014(INL) traz em seu artigo 3º, alínea "e", a definição dos agentes envolvidos:

[...] como operador final ou de frontend, "qualquer pessoa singular ou coletiva que exerça um grau de controle sobre um risco relacionado com a operação e o funcionamento do sistema de IA e que beneficie da sua operação." Por seu turno, a alínea "f" traz como operador inicial ou de backend "qualquer pessoa singular ou coletiva que, de forma contínua, defina as características da tecnologia, forneça dados e preste serviços essenciais de apoio de backend e, por conseguinte, exerça igualmente algum controle sobre o risco ligado à operação e ao funcionamento do sistema de IA.³⁰⁸

Como se pode perceber, a definição de operador perpassa dois elementos básicos: benefício e controle. Além disso, deve-se levar em consideração a complexidade dos sistemas de inteligência artificial, e, nesse aspecto Teresa Rodríguez De Las Heras Ballell apresenta cinco características que distinguem a Inteligência artificial das demais tecnologias: "autonomia crescente, dependência de dados, complexidade, vulnerabilidade e opacidade".³⁰⁹

Deste modo, na ocorrência de um dano injusto causado por um agente artificial autônomo, duas questões se sobressaem: A primeira é relativa a quem será atribuído o dever de reparação por esse dano; a segunda, atinente ao fundamento da responsabilidade civil capaz de justificar tal imposição. A eles, soma-se ainda um elemento agravante: a ininteligibilidade dos outputs desses sistemas.

A questão da responsabilidade civil em relação aos danos causados pelos sistemas de inteligência artificial ainda não tem uma solução definida, sendo necessários muitos debates sobre o assunto. Contudo, Filipe Medon traz algumas reflexões interessantes e que merecem ser destacadas:

Em verdade, o maior erro que se pode cometer é procurar fornecer resposta única para a tormentosa indagação: "quem responde pelos danos causados pela Inteligência Artificial?" Isso porque, como em boa hora se tem afirmado na Europa, não pode haver um único regime, pois não há uma só IA. Necessário se faz, portanto, considerar, por exemplo, a tipologia e a autonomia em concreto da IA envolvida no dano. Uma mesma tipologia, como

³⁰⁸ PARLAMENTO EUROPEU. **Resolução do Parlamento Europeu, de 20 de outubro de 2020, sobre o ato legislativo sobre os serviços digitais e questões relacionadas com os direitos fundamentais (2020/2022(INI))**. Disponível em: <https://eur-lex.europa.eu/legal-content/PT/TXT/PDF/?uri=CELEX:52020IP0274&from=SV>. Acesso em: 09 fev. 2022.

³⁰⁹ BALLELL, Teresa Rodríguez De Las Heras. Teresa La inteligencia artificial en clave jurídica. Propuesta de conceptualización y esbozo de los retos regulatorios. Una mirada europea. **In: Revista de Ciencia de la Legislación**, n. 8, outubro de 2020, Buenos Aires: Universidad del Salvador.

é o caso dos carros autônomos, pode ter diversos graus de autonomia em relação ao condutor humano. Significa dizer que eventualmente pode haver diferentes regimes aplicáveis dentro de uma única tipologia, que ainda depende de outras peculiaridades que podem surgir da dinâmica do acidente: havia obrigação da presença de condutor humano responsivo na retaguarda? O sistema de direção autônomo foi acionado fora de circunstâncias permitidas pelo fabricante? O usuário realizou todas as atualizações de segurança no software?

Além disso, há que se considerar que diversos também são os sujeitos envolvidos. Como também se tem afirmado na Europa e já se afirmava na referida obra publicada em 2020, há que se atentar para a legislação aplicável a cada pessoa natural ou jurídica que se relaciona ao dano. Está-se diante de relação de consumo? Atrai-se o Código de Defesa do Consumidor e seu regime de responsabilidade de natureza objetiva aos fornecedores pelo fato do produto. Está-se diante de relação empresarial, como é o caso dos danos decorrentes da atuação de administradores que delegam decisões para mecanismos automatizados no seio de companhias? Aplica-se regime de natureza subjetiva, baseado no dever de diligência, como se tem defendido em doutrina especializada.³¹⁰

Este é um assunto que deve ser debatido com cautela, com profundidade nas características próprias dos sistemas de inteligência artificial, não é assunto para ser definido com urgência, tanto é que nem mesmo a União Europeia que se propõe desde o final da última década a ser um marco referencial para a regulação da IA no mundo, não conta ainda com uma normativa aprovada, especialmente em termos de Responsabilidade Civil.

A Resolução mais recente do Parlamento Europeu, que data de 20 de outubro de 2020 e que contém recomendações à Comissão sobre o regime de Responsabilidade Civil aplicável à Inteligência Artificial (2020/2014(INL) trouxe, de plano, três importantes conclusões para o estudo do tema: (i) a desnecessidade de se realizar uma revisão completa das normas de responsabilidade civil existentes; (ii) o reconhecimento da importância de se analisar a tipologia da Inteligência Artificial para a definição do regime de responsabilidade civil aplicável, tendo em vista que IAs distintas implicam riscos distintos; e, por fim, (iii) a rejeição, pelo menos no presente momento, da criação de uma personalidade jurídica própria aos sistemas comandados por Inteligência Artificial.³¹¹

³¹⁰ MEDON, Filipe. Danos causados por inteligência artificial e a reparação integral posta à prova: por que o Substitutivo ao PL 21 de 2020 deve ser alterado urgentemente? *In: Migalhas*, 2021. Disponível em: <https://www.migalhas.com.br/coluna/migalhas-de-responsabilidade-civil/351200/danos-causados-por-inteligencia-artificial-e-a-reparacao-posta-a-prova>. Acesso em: 09 fev. 2022.

³¹¹ PARLAMENTO EUROPEU. **Resolução do Parlamento Europeu, de 20 de outubro de 2020, sobre o ato legislativo sobre os serviços digitais e questões relacionadas com os direitos fundamentais (2020/2022(INI))**. Disponível em: <https://eur-lex.europa.eu/legal-content/PT/TXT/PDF/?uri=CELEX:52020IP0274&from=SV>. Acesso em: 09 fev. 2022.

CONCLUSÃO

Compreender que a tecnologia com uso de sistemas de inteligência artificial está cada dia mais presente no cotidiano das pessoas e vem crescendo com uma velocidade imensurável é crucial para que se possa pensar em soluções para o novo contexto social. Além dos inúmeros benefícios que esses sistemas trazem, os mesmos podem, também, gerar danos que causam diversos tipos de repercussões na sociedade, pleiteando assim, que os poderes Legislativo e Judiciário realizem demandas que visem banir ou ao menos diminuir tais entraves. É, portanto, de extrema necessidade que a ciência jurídica reforce âmbitos de se precaver para dar uma resposta que, ainda que não consiga resolver todos os problemas, possa, ao menos, encontrar caminhos para reduzir os impactos e consiga direcionar todos num mesmo sentido de maneira objetiva.

O uso de Sistemas de Inteligência Artificial, seja de maneira direta, em softwares, como buscadores de internet, streaming de música, filmes e jogos, seja através de Sistema de Inteligência Artificial na modalidade embarcada ou, como Robôs e carros autônomos, as aplicações da IA estão amplamente difundidas ao redor do mundo, o que sugere fortemente uma urgência voltada para a necessidade de regulação, através da construção de governança ética para os Sistemas de Inteligência Artificial.

Apesar do avanço que essa tecnologia vem alcançando, se mostrando cada dia mais presente no cotidiano das pessoas, ainda restam algumas preocupações que fundamentam a construção de uma governança ética, garantindo a mitigação de consequências danosas às pessoas no uso dessa tecnologia. Sendo assim, a construção de normas técnicas e legais, para instituir padrões de segurança e ética a serem seguidos no desenvolvimento de Sistemas de Inteligência Artificial, visa que a tecnologia além de benéfica seja segura.

Sendo assim, a definição de um Marco Legal sobre inteligência artificial ainda é um desafio a ser enfrentado. Diante da ausência de regulamentação sobre a implementação da tecnologia de inteligência artificial, a presente dissertação tem o objetivo de analisar a necessidade de se criar parâmetros éticos para uma IA de Confiança pautados nos Direitos Humanos, para assim, desenvolver uma regulamentação mais eficaz. Assim, o problema de pesquisa que buscou-se responder é sobre quais os parâmetros que a IA de Confiança deve adotar para o

respeito aos Direitos Humanos e os Direitos Fundamentais?

No primeiro capítulo abordou as questões mais técnicas do que vem a ser a inteligência artificial. No qual se observou que a inteligência artificial é um somatório de vários campos da ciência, tendo em vista as constantes mudanças sociais que ensejaram sua criação e utilização.

Com isso, foi apresentado alguns conceitos essenciais para compreender como funciona o aprendizado dessas máquinas e a possível autonomia que elas podem adquirir. Destaca-se a compreensão de *machine learning* é o que chamamos de aprendizado de máquina, sendo elas supervisionadas, não supervisionada e por reforço. Atualmente, o mais utilizado é o aprendizado supervisionado, que se refere a utilização de dados para treiná-los e a partir deles contém a resposta desejada, como por exemplo, marcar um e-mail como spam ou não spam. O aprendizado não supervisionado vem ganhando espaço no cotidiano das pessoas, esse se assemelha mais com a forma que os seres humanos aprendem, o algoritmo aprende a representar (ou agrupar) as entradas submetidas, segundo medidas de similaridade, exemplos de aplicações de aprendizado não supervisionados são sistemas de recomendação de filmes ou músicas. Já o aprendizado por reforço, ainda pouco utilizado em nossa sociedade, ela é útil para se aprender como agir ou se comportar por sinais de recompensa e punição. Temos como exemplo, um automóvel que dirige “sozinho”, que toma decisões dependendo do cenário que observa ao redor, recebendo recompensas negativas quando colide com o ambiente ou com outros veículos, e com repetidas etapas, aos poucos “aprenda” a contornar os obstáculos.

Dentre estes tipos de aprendizagens, atualmente o mecanismo mais assertivo e seguro é o aprendizado supervisionado, os outros dois tipos de aprendizado ainda possuem margens de erro, bem como o problema chamado de *black box* ou caixa preta, que se trata de caminhos escolhidos pela máquina que ficam obscuros, que se mostram como desconhecidos para as pessoas que o utiliza, essa margem de obscuridade tem sido um dos grandes obstáculos para que se possa achar as soluções para os problemas apresentados pela máquina.

A partir da compreensão de como funciona o aprendizado das máquinas, se fez necessário apresentar alguns modelos já implementados no sistema judiciário de máquinas que utilizam de inteligência artificial com a finalidade otimizar o tempo e trazer mais produtividade para os envolvidos. Diante dos resultados até então obtidos por essas máquinas, apesar de se obter um índice de acurácia muito alto, ainda existe

a preocupação com respostas pautadas em vieses discriminatórios ou que ferem os Direitos Humanos e Direitos Fundamentais.

Essas problemáticas e outras que surgem com o desenvolvimento de máquinas inteligentes carecem de regulamentação própria, apesar de que alguns passos já foram dados nessa direção. Sendo assim, em um contexto mundial alguns Países avançaram em direção ao Marco Legal da Inteligência Artificial, dentre os principais estão a China, o Brasil, Estados Unidos e União Europeia. Em junho de 2019, a China divulgou regras de governança, lançando 8 princípios que devem ser seguidos no desenvolvimento de inteligência artificial no país. O governo dos Estados Unidos publicou em janeiro de 2020 uma proposta de regulamentação da inteligência artificial, a qual apresenta dez princípios para as agências governamentais aderirem ao propor regulamentos de Inteligência Artificial para o setor privado. Dentro do contexto europeu, pode-se mencionar a Carta de Ética Europeia sobre o Uso da Inteligência Artificial em Sistemas Judiciais e seu Ambiente, elaborada pela Comissão Europeia para Eficiência da Justiça e adotada na 31ª sessão plenária em Estrasburgo, nos dias 3 e 4 de dezembro de 2018. Com base na Carta de Ética, em 2019, foi elaborado pelo grupo de peritos de alto nível sobre a inteligência artificial (GPAN IA), instituição criada pela Comissão Europeia, um guia com as orientações éticas para uma IA de confiança, que prevê pilares como supervisão humana, segurança e robustez técnica, privacidade e governança de dados, transparência e não-discriminação na construção de sistemas inteligentes. Nesse contexto mundial, a União Europeia vem se mostrando a grande pioneira na construção dessas regras, tendo como foco principal garantir que se tenha uma Inteligência Artificial confiável a partir dos parâmetros éticos, e que possa assegurar a prevalência dos Direitos Humanos e dos Direitos Fundamentais.

No Brasil alguns passos também já foram dados, como propostas do Senado Federal existe o Projeto de Lei 5.051/2019, de 16 de setembro de 2019, do senador Styvenson Valentim (PODEMOS/RN), que estabelece os princípios para o uso da Inteligência Artificial no Brasil. O Projeto de Lei 5.691/2019, de 25 de outubro de 2019, também do senador Styvenson, que institui a Política Nacional de Inteligência Artificial. E, o Projeto de Lei 872/2021, de 12 março de 2021, do senador Veneziano Rêgo (MDB – PB), que dispõe sobre o uso da Inteligência Artificial no Brasil.

Recentemente aprovado pela Câmara dos Deputados, o Projeto de Lei n. 21/2020, de 04 de fevereiro de 2020, do deputado Eduardo Bismarck (PDT – CE), que tem o objetivo

de estabelecer princípios, direitos e deveres para o uso de Inteligência Artificial no Brasil, apresentando um conteúdo mais completo que as propostas do Senado, contudo, a legislação não é suficientemente profunda e se espera que o Senado Federal, no qual o projeto encontra-se em tramitação, melhore essas questões. Apesar de alguns avanços já terem sido direcionados para uma regulamentação do uso da Inteligência Artificial no Brasil, enquanto não há legislação aprovada, a melhor recomendação é observar práticas internacionais.

Portanto, para que se tenha uma IA de confiança é necessário que seu desenvolvimento se construa a partir da Ética, bem como, sejam reconhecidos os Direitos Humanos e os Direitos Fundamentais. Sendo assim, no terceiro capítulo, apresentou-se o conceito de Direitos Humanos proposto por Joaquin Herrera Flores, chamado de Teoria Crítica dos Direitos Humanos, o qual reconhece a complexidade e pluralidade social. O autor critica a ideia de universalizar uma realidade e uma noção de mundo como algo certo, como uma visão de mundo única, tendo em vista que esta percepção é oriunda do que hoje conhecemos como o mundo Ocidental.

No entanto, diante dessa visão crítica apresentada por Herrera Flores é necessário que se considere quatro condições. A primeira condição aborda a importância do entendimento concreto e existente da sociedade, sobretudo de suas imprecisões e falhas, para chegar numa concepção de realidade onde todos estão inseridos, bem como onde se quer chegar para alcançar o objetivo principal, que é a dignidade humana. Ainda, o autor salienta que a segunda condição é o reconhecimento de que o pensamento crítico é um pensamento de combate, por isso, carece de compreensão ao identificar o adversário e firmar os objetivos e as intenções nas lutas, para que estas sejam inseridas na sociedade e haja mobilização. A terceira condição está no fato de admitir que os Direitos Humanos assumem o caráter de lutas pela igualdade e o empoderamento dos grupos mais desfavorecidos ao lutar por novas formas, mais igualitárias e generalizadoras, de acesso aos bens protegidos pelo direito. A quarta e última condição se refere ao reconhecimento de que ter um pensamento crítico não significa necessariamente negação de algo, mas sim, na capacidade de afirmar as fraquezas de uma ideia, de um argumento, de um raciocínio, inclusive dos nossos, quando não forem consistentes, tentando corrigi-los para reforçá-los.

Salienta-se, portanto, que se concebeu na presente pesquisa a Teoria Crítica dos Direitos Humanos, proposta por Herrera Flores, por esta trazer um olhar mais

realista da sociedade, reconhecendo as complexidades existentes nos diferentes povos e suas diversas culturas. Sendo assim, se faz necessário observar que se confunde a complexidade empírica e a complexidade jurídica, no passo que o empírico compete ao fato de ter direitos, enquanto o jurídico; consiste na ideia de o que/como deveríamos ter através do direito.

Além disto, a ética é outro fator de grande relevância para a construção de uma IA de confiança, por essa ser parte essencial para o reconhecimento de uma Inteligência Artificial de confiança, revelando-se imperioso a discussão sobre a relação existente entre direito e moral. Tendo em vista, que o direito, como apreciação do fato ou fenômeno social, está subordinado à Moral, bem como à imposição de parâmetros nem sempre perfeitamente racionalizados, os quais se revelam vinculados a pressupostos de senso comum, míticos, religiosos, políticos, etc.

Diante disso, buscou-se apresentar as Orientações Éticas para uma IA de Confiança, proposta pela Comissão Europeia em 2018 como parâmetro de regulamentação da IA, por ter o foco voltado na construção ética, bem como estar inserida nos valores dos Direitos Humanos e os Direitos Fundamentais.

O uso da IA deve ser harmônico com os direitos garantidos aos cidadãos, sejam aqueles consagrados na constituição nacional ou pela Declaração Universal de Direitos Humanos. A utilização de novas tecnologias pelo mercado ou pelos Estados devem sempre estar alinhadas com o bem-estar social e com a promoção de uma vida digna, para isto, é necessário que se construa uma legislação bom base em parâmetros de uma IA de Confiança, tendo em vista que a confiança e a ética são os pontos centrais para a aplicação da inteligência artificial.

Segundo o documento, os sistemas de IA devem ser legal, ética e robusta, de modo a evitar erros ou a terem condição de lidar com estes, corrigindo eventuais inconsistências. Esses problemas podem ter sérios impactos na sociedade, como a discriminação de pessoas no acesso a um serviço ou até mesmo quedas de bolsas cujas compras e vendas de ações utilizam essas tecnologias.

Sendo assim, a construção de uma Inteligência Artificial deve essencialmente garantir o respeito aos Direitos Humanos, bem como a garantia dos Direitos Fundamentais previstos na Constituição Federal Brasileira de 1988. Com base nesta perspectiva observa-se que a Ética é caminho essencial para a construção de uma IA de confiança (termo utilizado pela Comissão Europeia ao criar o documento denominado “Orientações Éticas para uma IA de Confiança”).

E ao tratar da compreensão dos Direitos Humanos, não se pode esquecer que em um mundo globalizado, a IA serve para todo tipo de cultura social, para todo tipo de governo, inclusive para aqueles que não se estruturam em valores relacionados a direitos humanos, o mundo virtual não possui fronteiras. Por isso, se deu relevância a teoria crítica dos Direitos Humanos de Joaquín Herrera Flores, por ser desenvolvida com intenção de romper com o pensamento hegemônico dos direitos universais fundamentados nas noções ocidentais e eurocêtricas de dignidade humana, que alicerça o sistema neoliberal intensificando as desigualdades sociais. Desta forma, ao se pensar em uma regulamentação para a IA, mesmo que seja uma regulamentação interna, se faz necessário ter essa compreensão global.

No Brasil, atualmente vem ganhando destaque o Projeto de Lei n. 21/2020, aprovado pela Câmara dos Deputados no dia 29 de setembro de 2021 e encaminhado para o Senado Federal para votação. O Projeto, apensar e trazer princípios e fundamentos pautados na ética e os Direitos Fundamentais, ainda apresenta muitas lacunas.

Diante do exposto, demonstrou-se que, com o avanço da tecnologia e o desenvolvimento de inteligências artificiais, as atividades realizadas no âmbito do Direito podem ser facilitadas com a implementação de máquinas a fim de se alcançar mais celeridade e eficiência, seja no âmbito da advocacia ou no âmbito do poder público. Contudo, esta atividade contém especificações que devem ser discutidas em esfera legislativa.

Além disso, ao abordar qualquer assunto relacionado a IA, não se pode deixar de salientar a necessidade de uma estratégia de governança transnacional, pois seu reflexo ultrapassa as fronteiras estabelecidas pelos Estados Nações, e a partir destes parâmetros construir as regulamentações internas observando as particularidades de cada país.

Por este motivo, visando combater as barreiras ainda existentes na construção da IA, o documento emitido pela Comissão Europeia vem se mostrando uma forte referência para uma IA de Confiança, tendo em vista a sua capacidade de articulação dos elementos jurídicos relacionados à complexidade encontrada no tema tratado.

É importante salientar, que o que se pretende é uma profunda reflexão acerca das implicações éticas, legais e sociais que a implementação das tecnologias de inteligência artificiais causam na sociedade de forma abrangente a partir da sua

implementação no judiciário.

Pois, não restam dúvidas de que as máquinas tendem a consolidação de um sistema único digital e de dados que visa uma maior facilidade e celeridade em um âmbito jurídico. Com isso, para ser possível a instalação de sistemas de inteligência artificial no poder judiciário brasileiro, assim como no âmbito privado, se faz essencial que haja uma confiança entre aqueles que estarão sujeitos a máquina.

REFERÊNCIAS

- ALCOFORADO, Fernando A. G. **Os benefícios e os riscos da singularidade tecnológica baseada na Superinteligência Artificial**. 2020. Disponível em: https://www.portalsaudenoar.com.br/wp-content/uploads/2020/12/OS_BENEFICIOS_E_OS_RISCOS_DA_SINGULARIDA.pdf. Acesso em: 05 set. 2021.
- ALEX, Robert. **Teoria dos direitos fundamentais**. Tradução de Virgílio Afonso da Silva. 2. ed., São Paulo: Malheiros, 2015.
- AMORIM, Fernanda Pacheco. **Respeita as Mina: Inteligência artificial e violência contra a mulher**. 1. ed., Florianópolis: EMais Editora, 2019.
- ANGWIN, Julia; LARSON, Jeff; MATTU, Surya; KIRCHNER, Lauren. Machine Bias: There's software used across the country to predict future criminals. And it's biased against blacks. *In: ProPublica*, 2016. Disponível em: <https://www.propublica.org/article/machine-bias-risk-assessments-in-criminal-sentencing>. Acesso em: 02 ago. 2021.
- ARISTÓTELES. **Ética a Nicômaco**. Livro II. Coleção Os Pensadores, São Paulo: Nova Cultural, 1996.
- AZEREDO, João Fábio Azevedo e. **Reflexos do emprego de sistemas de inteligência artificial nos contratos**. 2014. 221 fls. Dissertação (Mestrado em Direito) – Universidade de São Paulo, São Paulo, 2014. Disponível em: https://www.teses.usp.br/teses/disponiveis/2/2131/tde-12122014-150346/publico/Dissertacao_reflexos_inteligencia_artificial_contratos_reduzida.pdf. Acesso em: 24 jul. 2021.
- BALLELL, Teresa Rodríguez De Las Heras. Teresa La inteligencia artificial en clave jurídica. Propuesta de conceptualización y esbozo de los retos regulatorios. Una mirada europea. *In: Revista de Ciencia de la Legislación*, n. 8, outubro de 2020, Buenos Aires: Universidad del Salvador.
- BARCELLOS, Ana Paula de. Normatividade dos princípios e o princípio da dignidade da pessoa humana na Constituição de 1988. *In: Revista de Direito Administrativo*, Rio de Janeiro, vol. 221, p. 159-188, jul./set. 2000. Disponível em: <http://bibliotecadigital.fgv.br/ojs/index.php/rda/article/view/47588/45167>. Acesso em: 09 out. 2021.
- BATITUCCI, Luiz Anísio Vieira; MARTINS, Amilar Domingos Moreira. **Sistemas apoiados por Inteligência Artificial**: Superior Tribunal de Justiça. 2019. Disponível em: <http://www.jf.gov.br/cjf/corregedoria-da-justica-federal/centro-de-estudos-judiciarios-1/eventos/eventos-cej/2019/stj-apresentacao-enatic-junho-2019-2.pdf>. Acesso em: 04 ago. 2021.
- BECK, Ulrich. **Sociedade de risco: rumo a uma outra modernidade**. 2. ed., São Paulo: Editora 34, 2011.

BENGIO, Yoshua; COURVILLE, Aaron; VINCENT, Pascal. Representation learning: A review and new perspectives. *In: IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 35, n. 8, p. 1798–1828, 2013.

BENTES, Dorinethe dos Santos; CURVINO, Sarah Fernandes. O acesso à Justiça e a Inteligência Artificial: uma análise da Resolução n.º 332/2020 do Conselho Nacional de Justiça. *In: XI Congresso RECAJ-UFMG*, Belo Horizonte, 2020.

BENTO, Leondardo V. **Governança e Governabilidade na reforma do Estado:** entre eficiência e democratização. Barueri: Manole, 2003.

BIRD, Sarah; BAROCAS, Solon; CRAWFORD, Kate; DIAS, Fernando; WALLACH, Hanna. Exploring or Exploiting? Social and Ethical Implications of Autonomous Experimentation in AI. *In: Workshop on Fairness, Accountability, and Transparency in Machine Learning*. 2016. Disponível em: <https://www.microsoft.com/en-us/research/publication/exploring-or-exploiting-social-and-ethical-implications-of-autonomous-experimentation-in-ai/>. Acesso em: 31 jul. 2021.

BITTAR, Eduardo Carlos Bianca. **Curso de ética jurídica:** ética geral e profissional. 5. ed. rev., São Paulo: Saraiva, 2007.

BITTAR, Eduardo Carlos Bianca; ALMEIDA, Guilherme Assis de. **Curso de Filosofia do Direito**. 8. ed. rev. aum., São Paulo: Atlas, 2010.

BOEING, Daniel Henrique Arruda. **Ensinando um robô a julgar:** pragmática, discricionariedade e vieses no uso de aprendizado de máquina no judiciário. 2019. 84 fls. Monografia (Graduação em Direito) – Universidade Federal de Santa Catarina, Florianópolis, 2019.

BRASIL. [Constituição, 1988]. **Constituição da República Federativa do Brasil de 1988**. Disponível em: http://www.planalto.gov.br/ccivil_03/constituicao/constituicao.htm. Acesso em: 17 out. 2021.

BRASIL. **Lei nº 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: 13 jan. 2022.

BRASIL. **Portaria GM nº 4.617, de 6 de abril de 2021**. Institui a Estratégia Brasileira de Inteligência Artificial e seus eixos temáticos. Disponível em: https://www.in.gov.br/en/web/dou/-/portaria-gm-n-4.617-de-6-de-abril-de-2021-*313212172. Acesso em: 27 dez. 2021.

BRASIL. Câmara dos Deputados. **Projeto de Lei n. 21, de 2020**. Estabelece princípios, direitos e deveres para o uso de inteligência artificial no Brasil, e dá outras providências. Disponível em: https://www.camara.leg.br/proposicoesWeb/prop_mostrarintegra;jsessionid=node015en4zcdsiu5ggqchdxtw9qjv11397236.node0?codteor=1853928&filename=PL+21/2020. Acesso em: 08 set. 2021.

BRASIL. Câmara dos Deputados. **Projeto de Lei n. 240, de 2020**. Cria a Lei da Inteligência Artificial, e dá outras providências. Disponível em: https://www.camara.leg.br/proposicoesWeb/prop_mostrarintegra?codteor=1857143&filename=PL+240/2020. Acesso em: 08 set. 2021.

BRASIL. Conselho Nacional de Justiça. **Resolução nº 332 de 21 de agosto de 2020**. Dispõe sobre a ética, a transparência e a governança na produção e no uso de Inteligência Artificial no Poder Judiciário e dá outras providências. Disponível em: <https://atos.cnj.jus.br/atos/detalhar/3429>. Acesso em: 08 set. 2021.

BRASIL. Ministério da Ciência, Tecnologia e Inovações. **Estratégia Brasileira de Inteligência Artificial (EBIA)**. 2021. Disponível em: https://www.gov.br/mcti/pt-br/acompanhe-o-mcti/transformacaodigital/arquivosinteligenciaartificial/ia_estrategia_documento_referencia_4-979_2021.pdf. Acesso em: 09 jan. 2022.

BRASIL. Senado Federal. **Projeto de Lei n. 5051, de 2019**. Estabelece os princípios para o uso da Inteligência Artificial no Brasil. Disponível em: <https://legis.senado.leg.br/sdleg-getter/documento?dm=8009064&ts=1630421610171&disposition=inline>. Acesso em: 08 set. 2021.

BRASIL. Senado Federal. **Projeto de Lei n. 5691, de 2019**. Institui a Política Nacional de Inteligência Artificial. Disponível em: <https://legis.senado.leg.br/sdleg-getter/documento?dm=8031122&ts=1630421785830&disposition=inline>. Acesso em: 08 set. 2021.

BRASIL. Senado Federal. **Projeto de Lei n. 872, de 2021**. Dispõe sobre o uso da Inteligência Artificial. Disponível em: <https://legis.senado.leg.br/sdleg-getter/documento?dm=8940096&ts=1634324707816&disposition=inline>. Acesso em: 08 set. 2021.

BUCHANAN, Bonnie G. **Artificial intelligence in Finance**. Washington: Zenodo, 2019. Disponível em: <https://zenodo.org/record/2612537#.YUaPM7hKjIV>. Acesso em: 28 ago. 2021.

CARBALLIDO, Manuel E. Gándara. **Los derechos humanos en el siglo XXI: una mirada desde el pensamiento crítico**. Buenos Aires: CLACSO, 2019.

CARLONI, Giovanna Louise Bodin de Saint-Ange Comnène. **Privacidade e Inovação na Era do Big Data**. 2013, 66 fls. Trabalho de Conclusão de Curso (Graduação em Direito) – Fundação Getúlio Vargas, Rio de Janeiro, 2013.

CASA Branca divulga princípios para a regulamentação de Inteligência Artificial. **In: Internetlab**, 2020. Disponível em: <https://www.internetlab.org.br/pt/itens-semanario/eua-casa-branca-divulga-principios-para-a-regulamentacao-de-inteligencia-artificial/>. Acesso em: 22 out. 2021.

CHINA: Lançados Princípios de Governança de IA. **In: Library of Congress**, 2019.

Disponível em: <https://www.loc.gov/item/global-legal-monitor/2019-09-09/china-ai-governance-principles-released/>. Acesso em: 22 out. 2021.

CLAVERA, Walter V. Aprendizado de Máquina (Machine Learning). *In: Sejatech e Falettech*, 2019. Disponível em: <https://www.redesdaude.com.br/aprendizado-de-maquina-machine-learning/>. Acesso em: 28 ago. 2021.

COELHO, João Victor de Assis Brasil Ribeiro. **Aplicações e Implicações da Inteligência Artificial no Direito**. 2017. 61 f. TCC (Graduação em Direito) – Universidade de Brasília, Brasília, 2017. Disponível em: https://bdm.unb.br/bitstream/10483/18844/1/2017_JoaoVictordeAssisBrasilRibeiroCoelho.pdf. Acesso em: 02 ago. 2021.

COELHO, Mateus. Inteligência Artificial e Deep Learning. *In: Laboratório iMobilis*, 2019. Disponível em: <http://www2.decom.ufop.br/imobilis/inteligencia-artificial-e-deep-learning/>. Acesso em: 30 jul. 2021.

COMISSÃO EUROPEIA. **Orientações éticas para uma IA de Confiança**. Junho de 2018. Disponível em: <http://escola.mpu.mp.br/servicos-academicos/atividades-academicas/inovaescola/curadoria/3-ciclo-de-debates/inteligencia-artificial-e-internet-das-coisas-oportunidades-e-desafios/ethicsguidelinesfortrustworthyai-ptpdf.pdf>. Acesso em: 20 out. 2021.

COMISSÃO EUROPEIA. **Proposta de regulamento que estabelece regras harmonizadas em matéria de Inteligência Artificial (Regulamento Inteligência Artificial) e altera determinados atos legislativos da União**. 21 abr. 2021. Disponível em: https://eur-lex.europa.eu/resource.html?uri=cellar:e0649735-a372-11eb-9585-01aa75ed71a1.0004.02/DOC_1&format=PDF. Acesso em: 08 set. 2021.

COMPARATO, Fábio Konder. **Ética, Direito, Moral e Religião no Mundo Moderno**. São Paulo: Companhia das Letras, 2006.

CONSELHO EUROPEU. European Commission for the Efficiency of Justice (CEPEJ). **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-december-2018/16808f699c>. Acesso em: 08 set. 2021.

CONSELHO NACIONAL DE JUSTIÇA. **Inteligência artificial: Trabalho judicial de 40 minutos pode ser feito em 5 segundos**. 2018. Disponível em: <https://www.cnj.jus.br/inteligencia-artificial-trabalho-judicial-de-40-minutos-pode-ser-feito-em-5-segundos/>. Acesso em: 28 ago. 2021.

CONSELHO NACIONAL DE JUSTIÇA. **Relatório Justiça em Números 2018: ano-base 2017**. 2018. Disponível em: www.cnj.jus.br/files/conteudo/arquivo/2018/09/8d9faee7812d35a58cee3d92d2df2f25.pdf. Acesso em: 29 ago. 2021.

COSSETTI, Melissa Cruz. **O que é inteligência artificial?** 2018. Disponível em: <https://tecnoblog.net/263808/o-que-e-inteligencia-artificial/>. Acesso em: 30 jul. 2021.

COSTA, Marcos Bemquerer; BASTOS, Patrícia Reis Leitão. Alice, Monica, Adele, Sofia, Carina e Ágata: o uso da inteligência artificial pelo Tribunal de Contas da União. **In: Revista do Tribunal de Contas do Estado de Goiás**, Belo Horizonte, ano 2, n. 3, p. 11-34, jan./jun. 2020.

COSTA, Suzana Rita da. **A Contribuição da Inteligência Artificial na Celeridade dos Trabalhos Repetitivos no Sistema Jurídico**. 2020. 72 fls. Dissertação (Mestrado em Mídia e Tecnologia) – Universidade Estadual Paulista Júlio de Mesquita Filho – Unesp, Bauru, 2020.

COUTO, João Victor Lima de Abreu. Desafios da inteligência artificial na interpretação jurídica em loomis vs. Wisconsin: novas tecnologias na interpretação de decisões judiciais sob uma ótica da hermenêutica de hans george gadamer. **In: Congresso Internacional de Direito e Inteligência Artificial - Inteligência artificial e tecnologias aplicadas ao direito II**. Belo Horizonte: Skema Business School, 2020.

DESORDI, Danubia; BONA, Carla Della. A inteligência artificial e a eficiência na administração pública. **In: Revista de Direito**, Viçosa, vol. 02, n. 02, 2020. Disponível em: <https://periodicos.ufv.br/revistadir/article/view/9112/5928>. Acesso em: 29 ago. 2021.

DIAS, Guilherme Ataíde; VIEIRA, Américo Augusto Nogueira. Big Data: questões éticas e legais emergentes. **In: Ciências da Informação**, Brasília, vol. 42, n. 2, p. 174-184, 2013. Disponível em: <http://revista.ibict.br/ciinf/article/view/1380/1558>. Acesso em: 24 jul. 2021.

DONEDA, Danilo Cesar Maganhoto; MENDES, Laura Schertel; SOUZA, Carlos Affonso Pereira de; ANDRADE, Norberto Nuno Gomes de. Considerações iniciais sobre inteligência artificial, ética e autonomia pessoal. **In: Pensar - Revista de Ciências Jurídicas**, Fortaleza, vol. 23, n. 4, p. 1-17, out./dez. 2018.

DORNELLE, João Ricardo. **O que são direitos humanos**. São Paulo. 5. ed., São Paulo: Editora Brasiliense Ltda, 2013.

FACELI, K.; LORENA, A. C.; GAMA, J.; CARVALHO, A. C. P. L. F. **Inteligência Artificial: Uma Abordagem de Aprendizagem de Máquina**. Rio de Janeiro: LTC, 2011.

FERNANDES, R. Dr^a Luzia, primeira robô-advogada do Brasil, já tem trabalho pela frente. **In: Instituto Humanista Unisinos**, 2017. Disponível em: <http://www.ihu.unisinos.br/78-noticias/569427-dr-luzia-primeira-robo-advogada-do-brasiljatem-trabalho-pela-frente>. Acesso em: 28 ago. 2021.

FERRAZ JUNIOR, Tércio Sampaio. **Introdução ao estudo do direito: técnica, decisão, dominação**. 2. ed., São Paulo: Atlas, 1994.

FERREIRA, Danielle Parra. **Análise Comparativa do uso do Big data em redes sociais**: Um estudo sobre o posicionamento de juízes e advogados. 2016. 126 fls.

Trabalho de Curso (Curso de Gestão da Informação) – Universidade Federal do Paraná, Curitiba, 2016.

FLORES, Joaquín Herrera. **Teoria crítica dos direitos humanos: os direitos humanos como produtos culturais**. Rio de Janeiro: Editora Lumen Juris, 2009.

FLORES, Joaquín Herrera. **A (re)invenção dos direitos humanos**. Florianópolis: Fundação Boiteux, 2009.

GALLARDO, Helio. Teoría crítica y derechos humanos: una lectura latinoamericana. **In: Revista de Derechos Humanos y Estudios Sociales**, São Luís Potosí, vol. 2, n. 4, p. 57-89, jul. 2010. Disponível em: <http://www.derecho.uaslp.mx/Documents/Revista%20REDHES/N%C3%BAmero%204/Redhes4-03.pdf>. Acesso em: 30 set. 2021.

GIBNEY, Elizabeth. Google AI algorithm masters ancient game of Go. **In: NATURE. International Weekly Journal of Science**, 2016. Disponível em: <http://www.nature.com/news/google-ai-algorithm-masters-ancient-game-of-go1.19234>. Acesso em: 26 jul. 2021.

GOMES, Helton Simões. Como as robôs Alice, Sofia e Monica ajudam o TCU a caçar irregularidades em licitações. **In: G1**, 18 mar. 2018. Disponível em: <https://g1.globo.com/economia/tecnologia/noticia/como-as-robos-alice-sofia-e-monica-ajudam-o-tcu-a-cacar-irregularidades-em-licitacoes.ghtml>. Acesso em: 29 ago. 2021.

GOMES, Rodrigo Dias de Pinho. **Desafios à Privacidade: Big data, consentimento, legítimos interesses e novas formas de legitimar o tratamento de dados pessoais**. 2017. Disponível em: <https://itsrio.org/wp-content/uploads/2017/03/Rodrigo-Gomes.doc-B.pdf>. Acesso em: 24 jul. 2021.

GOODFELLOW, Ian; BENGIO, Yoshua; COURVILLE, Aaron. **Deep Learning**. Disponível em: https://www.deeplearningbook.org/front_matter.pdf. Acesso em: 18 ago. 2021.

GUSMÃO, Paulo Dourado. **Introdução ao estudo do direito**. Rio de Janeiro: Forense, 1999.

HAYKIN, Simon. **Redes Neurais: princípios e prática**. 2. ed., Porto Alegre: Bookman Companhia, 2001.

IBM – International Business Machines Corporation. **10 formas como os dados podem ser sua vantagem estratégica**. 2018. Disponível em: <https://www.ibm.com/downloads/cas/MY9809DE>. Acesso em: 02 ago. 2021.

INAZAWA, Pedro; PEIXOTO, Fabiano Hartmann; CAMPOS, Teófilo de; SILVA, Nilton; BRAZ, Fabricio. **Projeto Victor: como o uso do aprendizado de máquina pode auxiliar a mais alta corte brasileira a aumentar a eficiência e a velocidade de avaliação judicial dos processos julgados**. 2019. Disponível em: https://cic.unb.br/~teodecampos/ViP/inazawa_etal_compBrasil2019.pdf. Acesso em:

02 ago. 2021.

KELSEN, Hans. **O que é justiça?:** a justiça, o direito e a política no espelho da ciência. Tradução de Luís Carlos Borges. São Paulo: Martins Fontes, 1998.

KNIGHT, Will. The Dark Secret at the Heart of AI. **In: MIT Technology Review**, 2017. Disponível em: <https://www.technologyreview.com/2017/04/11/51113/the-dark-secret-at-the-heart-of-ai/>. Acesso em: 30 jul. 2021.

KOLANOVIC, Marko; KRISHNAMACHARI, Rajesh T. **Big Data and AI Strategies: Machine Learning and Alternative Data Approach to Investing**. 2017. Disponível em: <https://cpb-us-e2.wpmucdn.com/faculty.sites.uci.edu/dist/2/51/files/2018/05/JPM-2017-MachineLearningInvestments.pdf>. Acesso em: 28 ago. 2021.

LAGE, Fernanda de Carvalho. **Manual de Inteligência Artificial no Direito Brasileiro**. Salvador: Editora Juspodivm, 2021.

LAGE, Lorena Muniz e Castro; LIMA, Henrique Cunha Souza; MARTINO, Antonio Anselmo. Inteligência Artificial e Tecnologias Aplicadas ao Direito – II. **In: Congresso Internacional de Direito e Inteligência Artificial**. Belo Horizonte: Skema Business School, 2020. Disponível em: <https://www.conpedi.org.br/wp-content/uploads/2020/09/SKEMA-Intelig%C3%Aancia-Artificial-e-tecnologias-aplicadas-ao-Direito-II.pdf>. Acesso em: 30 jul. 2021.

LANNES, Yuri Nathan da Costa; VALENTINI, Rômulo Soares; PIMENTA, Raquel Betty de Castro. Inteligência Artificial e Tecnologias Aplicadas ao Direito – III. **In: Congresso Internacional de Direito e Inteligência Artificial**. Belo Horizonte: Skema Business School, 2020. Disponível em: <https://www.conpedi.org.br/wp-content/uploads/2020/09/SKEMA-Intelig%C3%Aancia-Artificial-e-tecnologias-aplicadas-ao-Direito-III.pdf>. Acesso em: 02 ago. 2021.

LAWGEEX. **Comparing the Performance of Artificial Intelligence to Human Lawyers in the Review of Standard Business Contracts**. 2018. Disponível em: <https://images.law.com/contrib/content/uploads/documents/397/5408/lawgeex.pdf>. Acesso em: 29 ago. 2021.

LEGAL SaaS AI Platform LawGeex levanta US \$ 7 milhões na rodada de financiamento. **In: CISION PRNewswire**, 2017. Disponível em: <https://www.prnewswire.com/news-releases/legal-saas-ai-platform-lawgeex-raises-7-million-in-funding-round-615570484.html>. Acesso em: 29 ago. 2021.

LI, Deng; LIU, Yang. **Deep learning in natural language processing**. Suíça: Springer, 2018.

LIMA, Welton Dias de. Computadores e Inteligência: Uma explicação elucidativa sobre o Teste de Turing. **In: Revista Outras Palavras**, vol. 13, n. 1, 2017. Disponível em: <http://revista.faculdadeprojecao.edu.br/index.php/Projecao5/article/download/756/730>. Acesso em: 24 jul. 2021.

LÍVIO, Marcus. Um retrato sobre o uso da inteligência artificial no Poder Judiciário. *In: Justiça & Cidadania*. 2021. Disponível em: <https://editorajc.com.br/um-retrato-sobre-o-uso-da-inteligencia-artificial-no-poder-judiciario/>. Acesso em: 08 set. 2021.

LORENA, Ana Carolina; CARVALHO, André C. P. L. F. de. Uma introdução às Support Vector Machines. *In: Revista de Informática Teórica e Aplicada*. vol. 14, n.2, 2007. Disponível em: https://seer.ufrgs.br/rita/article/download/rita_v14_n2_p43-67/3543. Acesso em: 28 ago. 2021.

LUDWIG, Oswaldo Jr.; MONTGOMERY, Eduard. **Redes Neurais: Fundamentos e Aplicações em C**. 1. ed., Rio de Janeiro: Editora Ciência Moderna, 2007.

MAGRANI, Eduardo. **Entre dados e robôs: ética e privacidade na era da hiperconectividade**. 2. ed., Porto Alegre: Arquipélogo Editorial, 2019.

MAINI, Vishal; SABRI, Samer. **Machine Learning for Humans**. 2019. Disponível em: <https://everythingcomputerscience.com/books/Machine%20Learning%20for%20Humans.pdf>. Acesso em: 08 set. 2021.

MANICA, Rafael. Aplicação de uma rede neural artificial simplificada para a identificação de gradação de depósitos turbiditico. *In: Geociências*, vol. 32, n. 3, p. 429-440, 2013. Disponível em: https://www.researchgate.net/publication/260201326_APLICACAO_DE_UMA_REDE_NEURAL_ARTIFICIAL_SIMPLIFICADA_PARA_A_IDENTIFICACAO_DE_GRADACAO_DE_DEPOSITOS_TURBIDITICO. Acesso em: 29 ago. 2021.

MATTHIAS, Andreas. The responsibility gap: ascribing responsibility for the actions of learning automata. *In: Ethics and Information Technology*, v. 6, n. 3, p. 175-183, 2004.

MAYBIN, Simon. Sistema de algoritmo que determina pena de condenados cria polêmica nos EUA. *In: BBC News*. 2016. Disponível em: <https://www.bbc.com/portuguese/brasil-37677421>. Acesso em: 30 jul. 2021.

MCCARTHY, John. **What is Artificial Intelligence?** Califórnia, Stanford University, 2007. Disponível em: <http://jmc.stanford.edu/articles/whatisai/whatisai.pdf>. Acesso em: 26 jul. 2021.

MCCORDUCK, Pamela. **Machines who think: a personal inquiry into the history and prospects of artificial intelligence**. 2. ed., Massachusetts: A K Peters, 2004.

MEDON, Filipe. Danos causados por inteligência artificial e a reparação integral posta à prova: por que o Substitutivo ao PL 21 de 2020 deve ser alterado urgentemente? *In: Migalhas*, 2021. Disponível em: <https://www.migalhas.com.br/coluna/migalhas-de-responsabilidade-civil/351200/danos-causados-por-inteligencia-artificial-e-a-reparacao-posta-a-prova>. Acesso em: 09 fev. 2022.

MELO, João Ozório de. **Escritório de Advocacia estreia primeiro “robô-**

advogado” nos EUA. 2016. Disponível em: <https://www.conjur.com.br/2016-mai-16/escritorio-advocaciaestreiaprimeiro-robo-advogado-eua>. Acesso em: 02 ago. 2021.

MOGENDORFF, Janine Regina. A Escola de Frankfurt e seu legado. *In: Verso e Reverso*, Porto Alegre, vol. 26, n. 63, p. 152-159, set. 2012.

MONARD, Maria Carolina; BARANAUSKAS, José Augusto. Aplicações de Inteligência Artificial: Uma Visão Geral. *In: Anais – Congresso de Lógica Aplicada à Tecnologia*, São Paulo: Faculdade SENAC de Ciências Exatas e Tecnologia, 2000.

MURPHY, Kevin P. **Machine Learning: A Probabilistic Perspective.** 2012. Disponível em: <https://www.cs.ubc.ca/~murphyk/MLbook/pml-intro-22may12.pdf>. Acesso em: 28 ago. 2021.

NAVES, Fernanda de Moura Ribeiro. **Relatório de visita realizada ao Tribunal de Contas da União para exposição sobre os Sistemas Labcontas, Alice, Sofia e Mônica.** Goiânia, 2018. Disponível em: https://files.cercomp.ufg.br/weby/up/949/o/Relat%C3%B3rio_de_visita_ao_TCU_sobre_ALICE.pdf. Acesso em: 29 ago. 2021.

NUNES, Dierle; MARQUES, Ana Luiza Pinto Coelho. Inteligência Artificial e direito processual: vieses algorítmicos e os riscos de atribuição de função decisória às máquinas. *In: Revista de Processo*, São Paulo, vol. 285, p. 421-447, 2018. Disponível em: https://www.academia.edu/37764508/INTELIG%C3%8ANCIA_ARTIFICIAL_E_DIREITO_PROCESSUAL_VIESES_ALGOR%C3%8DTMICOS_E_%C3%83O_DE_FUN%C3%87%C3%83O_DECIS%C3%93RIA_%C3%80S_M%C3%81QUINAS_Artificial_intelligence_and_procedural_law_algorithmic_bias_and_the_risks_of_assignment_of_decision_making_function_to_machines. Acesso em: 30 jul. 2021.

O'NEIL, Cathy. **Weapons of math destruction.** New York: Broadway Books, 2016.

ORGANIZAÇÃO DAS NAÇÕES UNIDAS. **Declaração Universal dos Direitos Humanos.** 1948. Disponível em: <https://www.oas.org/dil/port/1948%20Declara%C3%A7%C3%A3o%20Universal%20dos%20Direitos%20Humanos.pdf>. Acesso em: 09 out. 2021.

PARENTONI, Leonardo Netto; VALENTINI, Rômulo Soares; ALVES, Tárík César Oliveira e. Panorama da regulação da inteligência artificial no Brasil: com ênfase no PLS n. 5.051/2019. *In: Revista Eletrônica do Curso de Direito da UFSM*, v. 15, n. 2 / 2020. Disponível em: <https://periodicos.ufsm.br/revistadireito/article/download/43730/pdf/246146>. Acesso em: 09 fev. 2022.

PARLAMENTO EUROPEU. **Resolução do Parlamento Europeu, de 20 de outubro de 2020, sobre o ato legislativo sobre os serviços digitais e questões relacionadas com os direitos fundamentais (2020/2022(INI)).** Disponível em:

<https://eur-lex.europa.eu/legal-content/PT/TXT/PDF/?uri=CELEX:52020IP0274&from=SV>. Acesso em: 09 fev. 2022.

PARLAMENTO Europeu aprova maior regulamentação de Inteligência Artificial na UE. **In: TSF rádio notícias**, 2020. Disponível em: <https://www.tsf.pt/futuro/parlamento-europeu-aprova-maior-regulamentacao-de-inteligencia-artificial-na-ue-12944972.html>. Acesso em: 13 jan. 2022.

PEIXOTO, Erick Lucena Campos; EHRHARDT JÚNIOR, Marcos. Breves notas sobre a resignificação da privacidade. **In: Revista Brasileira de Direito Civil – RBDCivil**, Belo Horizonte, v. 16, p. 35-56, abr./jun. 2018.

PEIXOTO, Fabiano Hartmann; SILVA, Roberta Zumblick Martins. **Inteligência Artificial e Direito**. Curitiba: Alteridade Editora, 2019.

PEREIRA, Jorge Luís. **Análise Preditiva em Sistemas de Informação no contexto do Big Data**. 2014. 72 fls. Trabalho de Conclusão de Curso (Curso de Sistemas de Informação) – Fundação de Ensino “Eurípes Soares da Rocha”, Marília, 2014.

PEREIRA, Thiago Pedroso. **A legalidade e efetividade dos atos judiciais realizados por inteligência artificial**. 2020. 122 fls. Dissertação (Mestrado em Direito) – Universidade Nove de Julho, São Paulo, 2020. Disponível em: <https://bibliotecatede.uninove.br/bitstream/tede/2411/2/Thiago%20Pedroso%20Pereira.pdf>. Acesso em: 27 dez. 2021.

PINHEIRO, Pe. José Ernanne Pinheiro; SOUSA JUNIOR, José Geraldo de; DINIS, Melillo et al. (Org.). **Ética, justiça e direito: reflexões sobre a reforma do judiciário**. 2. ed., Petrópolis: Vozes; CNBB, 1996,

PINTO, Henrique Alves. A utilização da inteligência artificial no processo de tomada de decisões. **In: RIL Brasília**, ano 57, n. 225, p. 43-60, jan./mar. 2020. Disponível em: https://www12.senado.leg.br/ril/edicoes/57/225/ril_v57_n225_p43.pdf. Acesso em: 30 jul. 2021.

PIOVESAN, Flávia. **Direitos humanos e o direito constitucional internacional**. 14. ed., São Paulo: Saraiva, 2013.

REALE, Miguel. **Filosofia do direito**. 19. ed., São Paulo: Saraiva, 1999.

REGULAMENTAÇÃO da inteligência artificial: entenda o debate mundial. **In: Alana**, [s.d.]. Disponível em: <https://blog.alana.ai/cresca-com-alana/regulamentacao-da-inteligencia-artificial-entenda-o-debate-mundial/>. Acesso em: 22 out. 2021.

RIBEIRO, Ana Lídia Lira. **Discriminação em algoritmos de inteligência artificial: uma análise acerca da Igpd como instrumento normativo mitigador de vieses discriminatórios**. 2021. 61 fls. Monografia (Graduação em Direito) – Universidade Federal do Ceará, Fortaleza, 2021.

RICHARDS, Neil M.; SMART, William D. **How should the law think about robots?**. 2016. Disponível em: <http://robots.law.miami.edu/wp->

content/uploads/2012/03/RichardsSmart_HowShouldTheLawThink.pdf. Acesso em: 30 jul. 2021.

ROCKCONTENT. **Saiba como funciona um algoritmo e conheça os principais exemplos existentes no mercado.** 2019. Disponível em: <https://rockcontent.com/br/blog/algoritmo/>. Acesso em: 5 set. 2021.

ROSS INTELLIGENCE. **Artificial Intelligence (AI) for the practice of law:** An introduction. 2018. Disponível em: <https://blog.rossintelligence.com/post/ai-introduction-law>. Acesso em: 02 ago. 2021.

RUSSELL, Stuart; NORVIG, Peter. **Artificial intelligence:** a modern approach. New Jersey: Editora Pearson Prentice Hall, 2010.

RUSSELL, Stuart; NORVIG, Peter. **Inteligência Artificial:** Uma abordagem moderna. 3. ed., Rio de Janeiro: Elsevier, 2013.

RUSSELL, Stuart; NORVIG, Peter. **Inteligência Artificial.** 3. ed., Rio de Janeiro: Editora Campus-Elsevier, 2013.

SALOMÃO, Luis Felipe (coord.). **Inteligência Artificial:** Tecnologia aplicada à gestão dos conflitos no âmbito do Poder Judiciário brasileiro. Fundação Getúlio Vargas, 2019. Disponível em: https://ciapj.fgv.br/sites/ciapj.fgv.br/files/estudos_e_pesquisas_ia_1afase.pdf. Acesso em: 04 ago. 2021.

SANTOS, Andréia. **O impacto do Big Data e dos Algoritmos nas Campanhas Eleitorais.** 2017. Disponível em: <https://itsrio.org/wp-content/uploads/2017/03/Andreia-Santos-V-revisado.pdf>. Acesso em: 24 jul. 2021.

SARAIVA, Sérgio. Deep Learning passo a passo. **In: Revista Programar.** 2018. Disponível em: <https://www.revista-programar.info/artigos/deep-learning-passo-passo/>. Acesso em: 08 set. 2021.

SARLET, Ingo. **Eficácia dos Direitos Fundamentais.** Porto Alegre: Ed. Livraria do Advogado, 2010.

SARLET, Ingo Wolfgang. **Dignidade (da pessoa) humana e direitos fundamentais na Constituição Federal de 1988.** 10. ed., Porto Alegre: Livraria do Advogado Editora, 2019.

SEARLE, John R.. Minds, Brains and Programs, **In: The Behavioral and Brain Sciences**, vol. 3, n. 3, p. 417-457, Cambridge: Cambridge University Press, 1980.

SERAPIÃO, Fábio. Dra. Luzia. **In: Estadão**, 13 de maio de 2018. Disponível em: <https://politica.estadao.com.br/blogs/fausto-macedo/dra-luzia/>. Acesso em: 29 ago. 2021.

SHIPPER, Laurianne-Marie. **Algoritmos que discriminam:** uma análise jurídica da discriminação no âmbito das decisões automatizadas e seus mitigadores. 2018.

57 fls. Monografia (Graduação em Direito) – Escola de Direito de São Paulo, Fundação Getúlio Vargas, São Paulo, 2018. Disponível em: <http://bibliotecadigital.fgv.br/dspace/bitstream/handle/10438/29878/Algoritmos%20que%20discriminam%20-%20Laurianne-Marie%20Schippers.pdf?sequence=1&isAllowed=y>. Acesso em: 05 set. 2021.

SILVA, Nilton Correia da. Notas iniciais sobre a evolução dos algoritmos do VICTOR: o primeiro projeto de inteligência artificial em supremas cortes do mundo. *In: FERNANDES, Ricardo Vieira de Carvalho; CARVALHO, Angelo Gamba Prata de (Coord.). Tecnologia jurídica & direito digital: II Congresso Internacional de Direito, Governo e Tecnologia*, p. 89-94, Belo Horizonte: Fórum, 2018.

STEIN, Raphael Wilson L. É Preciso Evolução Na Inteligência Artificial? *In: JusBrasil*. 2019. Disponível em: <https://stein.jusbrasil.com.br/artigos/743864200/e-preciso-evolucao-na-inteligencia-artificial>. Acesso em: 02 ago. 2021.

SUÍÇA vai testar robôs-carreiros a partir de setembro. *In: Veja*, 2016. Disponível em: <https://veja.abril.com.br/economia/suica-vai-testar-robos-carreiros-a-partir-de-setembro/>. Acesso em: 30 jul. 2021.

SUPERIOR TRIBUNAL DE JUSTIÇA. **Revolução tecnológica e desafios da pandemia marcaram gestão do ministro Noronha na presidência do STJ**. 2020. Disponível em: <https://www.stj.jus.br/sites/portalp/Paginas/Comunicacao/Noticias/23082020-Revolucao-tecnologica-e-desafios-da-pandemia-marcaram-gestao-do-ministro-Noronha-na-presidencia-do-STJ.aspx>. Acesso em: 04 ago. 2021.

SUPREMO TRIBUNAL FEDERAL. **Inteligência artificial vai agilizar a tramitação de processos no STF**. 2018. Disponível em: <http://www.stf.jus.br/portal/cms/verNoticiaDetalhe.asp?idConteudo=380038>. Acesso em: 10 ago. 2021.

SUTTON, Richard S.; BARTO, Andrew G. **Reinforcement learning: An introduction**. 2. ed., London: MIT press, 1998.

TAURION, César. **Big data**. Rio de Janeiro: Brasport Livros e Multimídia Ltda, 2013.

TOLEDO, Eduardo S. Projetos de inovação tecnológica na Administração Pública. *In: FERNANDES, Ricardo Vieira de Carvalho; CARVALHO, Angelo Gamba Prata de (Coord.). Tecnologia jurídica & direito digital: II Congresso Internacional de Direito, Governo e Tecnologia*. 2018. Belo Horizonte: Fórum, 2018.

TOSI, Giuseppe. Os Direitos Humanos: Reflexões iniciais. *In: TOSI, Giuseppe (coord.). Direitos Humanos: história, teoria e prática*. João Pessoa: Editora UFPB, 2004. Disponível em: https://www.academia.edu/40425908/Direitos_Humanos_Teoria_e_Pratica_Giuseppe_Tosi. Acesso em: 09 out. 2021.

TOSTES, Marcelo. Marco Legal da Inteligência Artificial traz benefícios para o Brasil. *In: Exame*, 2021. Disponível em: <https://exame.com/bussola/marco-legal-da->

inteligencia-artificial-traz-beneficios-para-o-brasil/. Acesso em: 13 set. 2021.

UNIÃO EUROPEIA. **Regulamento (UE) 2016/679 do Parlamento Europeu e do Conselho de 27 de abril de 2016.** Regulamento Geral sobre a Proteção de Dados. Disponível em: <https://eur-lex.europa.eu/legal-content/PT/TXT/?uri=celex%3A32016R0679>. Acesso em: 13 jan. 2022.

URWIN, Richard. **Artificial Intelligence: The Quest for the Ultimate Thinking Machine.** London: Arcturus, 2016.

VALENTINI, Romulo Soares. **Julgamento por computadores?** As novas possibilidades da juscibernética no século XXI e suas implicações para o futuro do direito e do trabalho dos juristas. 2017. 152 fls. Tese (Doutorado em direito) – Universidade Federal de Minas Gerais, Belo Horizonte, 2017.

VASCONCELLOS, Carlos Eduardo. “A legislação europeia sobre IA poderá ter influência no mundo todo”. *In: Consumidor Moderno*, 2021. Disponível em: <https://www.consumidormoderno.com.br/2021/07/02/legislacao-europeia-ia-influencia-mundo/>. Acesso em: 08 set. 2021.

VECCHI, Guilherme Busato; RISSO, Thábata Pontes Lazarou. **Solução do Desafio de Esquemas de Winograd em Português Brasileiro.** 2020. 94 fls. Monografia (Graduação em Engenharia) – Escola Politécnica da Universidade de São Paulo, São Paulo, 2020.

WINSTON, Patrick Henry. **Artificial Intelligence.** 3. ed., Boston: Addison-Wesley Publishing Company, 1993.