

Universidade Federal de São Carlos
Centro de Educação e Ciências Humanas
Programa de Pós-Graduação em Ciência, Tecnologia e Sociedade

**As diferenças nos processos de tomada de decisão da
inteligência artificial e dos seres humanos no
laboratório jogo de xadrez**

Léo Pasqualini de Andrade

São Carlos - SP
2026

LEO PASQUALINI DE ANDRADE

**As diferenças nos processos de tomada de decisão da
inteligência artificial e dos seres humanos no
laboratório jogo de xadrez**

Tese de Doutorado apresentada ao
Programa de Pós-Graduação em
Ciência, Tecnologia e Sociedade, do
Centro de Educação e Ciências
Humanas, da Universidade Federal de
São Carlos, como parte dos requisitos
para a obtenção do título de Doutor em
Ciência, Tecnologia e Sociedade.

Orientadora: Profa. Dra. Sylvia Iasulaitis
Coorientador: Prof. Dr. Alan Demétrius Baria Valejo

São Carlos – SP
2026

Andrade, Leo Pasqualini de

As diferenças nos processos de tomada de decisão da inteligência artificial e dos seres humanos no laboratório jogo de xadrez / Leo Pasqualini de Andrade -- 2026. 274f.

Tese (Doutorado) - Universidade Federal de São Carlos, campus São Carlos, São Carlos

Orientador (a): Sylvia Iasulaitis

Banca Examinadora: Luciana de Souza Gracioso, Roniberto Morato do Amaral, Eanes Torres Pereira,

Valério Brusamolin, Cláudio Alves de Amorim

Bibliografia

1. Inteligência artificial. 2. Cognição. 3. Xadrez. I.

Andrade, Leo Pasqualini de. II. Título.

Ficha catalográfica desenvolvida pela Secretaria Geral de Informática (SIn)

DADOS FORNECIDOS PELO AUTOR

Bibliotecário responsável: Arildo Martins - CRB/8 7180

Folha de Aprovação

Defesa de Tese de Doutorado do candidato Léo Pasqualini de Andrade, realizada em 23/04/2026.

Comissão Julgadora:

Profa. Dra. Sylvia lasulaitis (PPGCTS/UFSCar)

Profa. Dra. Luciana de Souza Gracioso (PPGCTS/UFSCar)

Prof. Dr. Roniberto Morato do Amaral (PPGCTS/UFSCar)

Prof. Dr. Eanes Torres Pereira (PPGCC/UFCG)

Prof. Dr. Valério Brusamolin (PPGCTS/IFPR)

Prof. Dr. Cláudio Alves de Amorim (MNPEF/UNEB)

DEDICATÓRIA

Este trabalho é dedicado ao meu pai, Léo (in memoriam);

À minha mãe Iracy e tia Jacy (in memoriam), gêmeas que completariam 200 anos em 2026

AGRADECIMENTOS

Agradeço muito à minha orientadora, Sylvia Iasulaitis, que resgatou para o doutorado e para o tema pelo qual me identifico desde a adolescência, me trouxe incentivo e teve paciência comigo. Devo muito a ela por eu ter conseguido realizar este doutorado.

Agradeço ao meu coorientador, Alan Demétrius Baria Valejo que me trouxe muitos conselhos e ajudou a clarear os pensamentos sobre a tese.

Agradeço ao Bruno Greco a importante ajuda com os testes com o ChatGPT.

Quero agradecer meus companheiros da UFSCar, o José Victor Catalano, pelas conversas e incentivos; Marcelo Seixas pelo apoio. Agradecimento especial à Vanessa Custódio pelo apoio que recebi no início do doutorado.

Meu agradecimento aos Grande Mestres Rafael Leitão e Alexandr Fier. Ao Aron Corrêa e Edmundo Aoyama pelos debates que me trouxeram muitas reflexões. Agradeço ao Luiz Paulo Rouanet pela ajuda e colaboração com os participantes da pesquisa. Agradeço o apoio que recebi da Jéssica Januário, Cristiana Fiusa, Paulo Nicolini, Cleiton Santana, enxadristas-pesquisadores e ao Mestre Internacional Cícero Braga pela ajuda com os livros.

Aos professores Marcos Castilho, Bruno Müller, Michelle Schultz, Ítalo Venturelli, meus agradecimentos especiais e gratidão. Agradecimento especial ao Eanes Pereira, Roberto Ferrari, Luciana Gracioso e Roniberto Amaral. Ao Cláudio Amorim que me ajudou muito em conversas epistemológicas sobre xadrez e inteligência artificial.

Quero deixar um agradecimento especial a minha mãe Iracy que me aturou bastante e a minha tia Jacy (in memoriam), ao meu irmão Marco, Ivan e Silvana.

Agradeço ao meu filho Pedro que me ajudou com testes dos diagramas.

Ao meu amigo Vander Pereira, que me ajudou em debates e companheirismo neurocientífico.

Agradeço aos colegas de conversas “ao pé da tela”, Célio, Miguel, Helgis, Sérgio e Walther (in memoriam).

Agradeço a Marisa Bassi, que suportou as conversas sobre a tese e me incentivou em muitos momentos de desânimo.

Aos enxadristas em geral que colaboraram com a tese. E a todos aqueles que direta ou indiretamente contribuíram para a realização deste trabalho e acreditam na Educação como instrumento de mudança social.

“CAMPBELL: ...Eisenhower entrou numa sala repleta de computadores e propôs às máquinas a seguinte questão: ‘Existe um Deus?’ Todas começam a funcionar, luzes se acendem, carretéis giram e após algum tempo uma voz diz: ‘Agora existe’” (Campbell e Moyers, 1991, p. 31).

“Oh no! I am losing!” thought the courtier in horror. “How is this possible? How can a mere machine, a combination of pulleys and wells, best me, the creation of God?” (Müller e Schaeffer, 2018, p. 12)

RESUMO

A Inteligência Artificial é um campo de estudos interdisciplinar. Os pioneiros da ciência da computação, matemáticos e engenheiros, vislumbraram a possibilidade de um computador simular o pensamento humano através de diferentes métodos computacionais, por exemplo, redes neurais artificiais, inspirados nas pesquisas biológicas da então embrionária neurociência. Filósofos e psicólogos se juntaram às pesquisas do campo, abordando questões sobre a possibilidade de a inteligência artificial pensar e ter consciência. O laboratório de testes para a inteligência artificial foi o xadrez, jogo milenar, com regras simples e de informação completa, que, para se jogar, desvela elementos do raciocínio humano, por exemplo, atenção seletiva, reconhecer padrões, abstrações e raciocínio lógico. Os cálculos de variantes no xadrez são realizados com muita acurácia pelos atuais algoritmos de inteligência artificial como Stockfish e AlphaZero, mas ainda assim apresentam questões a serem resolvidas pelos cientistas da computação a fim de alcançar soluções concebidas pela abstração humana. Embora o desempenho das *engines* de xadrez tenha sido amplamente estudado sob a perspectiva computacional, há uma lacuna na literatura quanto à comparação sistemática entre os processos de tomada de decisão humanos e algorítmicos no mesmo ambiente experimental. O objetivo desta tese é conhecer as diferenças, simetrias e controvérsias entre a inteligência artificial e a inteligência humana nas tomadas de decisões a partir do laboratório jogo de xadrez. Esta tese contribui para a literatura ao oferecer uma análise comparativa interdisciplinar inédita dos processos de tomada de decisão humanos e algorítmicos no laboratório do jogo de xadrez, articulando evidências empíricas, fundamentos cognitivos e interpretação sociotécnica. Os passos metodológicos utilizados se basearam na Teoria Ator-Rede, de Bruno Latour, que propõe a ideia de actantes, onde humanos e não humanos se juntam em uma rede de ações. Foi realizada uma revisão bibliográfica com as palavras-chave xadrez, inteligência e tomadas de decisão; um *survey* online com enxadristas sobre aspectos de reconhecimento de padrões, abstração e para se obter opiniões de situações-problema no contexto do xadrez. Também foram realizadas entrevistas semiestruturadas com cientistas da computação, neurocientistas e mestres de xadrez a fim de discutir o tema sobre inteligência artificial e o jogo de xadrez. Os resultados indicaram que, enquanto as *engines* de xadrez computam predominantemente por cálculo probabilístico e busca em árvore, os humanos mobilizam reconhecimento de padrões de sua memória de longo prazo adquiridos através de estudos e aprendizados, e de componentes emocionais e contextuais. Observou-se que enxadristas com maior força de jogo apresentaram maior capacidade de reconhecer padrões, enquanto jogadores menos experientes recorreram a cálculos mentais. As entrevistas revelaram controvérsias entre cientistas da computação, neurocientistas e mestres de xadrez quanto à noção das *engines* simularem emoções e fazer abstrações. Conclusão: apesar dos algoritmos de xadrez já terem superado os humanos no jogo, ainda assim são constatadas controvérsias destes programas em entender situações da partida, de não considerarem o viés artístico-cultural, ético e emocional, características estas associadas ao comportamento e à criatividade humanas. Por outro lado, os programas de xadrez são ferramentas que auxiliam os mestres em suas análises antes e depois das partidas jogadas e colaboram com jogadores de todos os níveis como fonte de lazer e aprendizagem.

Palavras-chave: Inteligência artificial; Tomadas de decisão. Xadrez. Cognição.

ABSTRACT

Artificial Intelligence is an interdisciplinary field of study. Pioneers in computer science, mathematicians, and engineers envisioned the possibility of a computer simulating human thought through different computational methods, for example, artificial neural networks, inspired by biological research in the then-embryonic neuroscience. Philosophers and psychologists joined the field's research, addressing questions about the possibility of artificial intelligence thinking and having consciousness. The testing ground for artificial intelligence was chess, an ancient game with simple rules and complete information, which, in order to play it, reveals elements of human reasoning, such as selective attention, pattern recognition, abstractions, and logical reasoning. The calculations of variations in chess are performed with great accuracy by current artificial intelligence algorithms such as Stockfish and AlphaZero, but still present questions to be resolved by computer scientists in order to reach solutions conceived by human abstraction. Although the performance of chess engines has been extensively studied from a computational perspective, there is a gap in the literature regarding the systematic comparison between human and algorithmic decision-making processes in the same experimental environment. The objective of this thesis is to understand the differences, symmetries, and controversies between artificial intelligence and human intelligence in decision-making, based on a chess game laboratory. This thesis contributes to the literature by offering a novel interdisciplinary comparative analysis of human and algorithmic decision-making processes in the chess game laboratory, articulating empirical evidence, cognitive foundations, and sociotechnical interpretation. The methodological steps used were based on Bruno Latour's Actor-Network Theory, which proposes the idea of actants, where humans and non-humans join together in a network of actions. A literature review was conducted using the keywords chess, intelligence, and decision-making; an online survey was carried out with chess players on aspects of pattern recognition, abstraction, and to obtain opinions on problem situations in the context of chess. Semi-structured interviews were also conducted with computer scientists, neuroscientists, and chess masters to discuss the topic of artificial intelligence and the game of chess. The results indicated that, while chess engines predominantly compute using probabilistic calculations and tree search, humans mobilize pattern recognition from their long-term memory acquired through study and learning, as well as emotional and contextual components. It was observed that chess players with greater playing strength showed a greater ability to recognize patterns, while less experienced players resorted to mental calculations. The interviews revealed controversies among computer scientists, neuroscientists, and chess masters regarding the notion of engines simulating emotions and making abstractions. Conclusion: although chess algorithms have already surpassed humans in the game, controversies are still observed regarding these programs' understanding of game situations and their failure to consider artistic-cultural, ethical, and emotional biases—characteristics associated with human behavior and creativity. On the other hand, chess programs are tools that assist masters in their analyses before and after games and help players of all levels as a source of leisure and learning.

Keywords: Artificial intelligence. Decision making. Chess. Cognition.

Lista de Figuras

Figura 1 - Os passos metodológicos da pesquisa	31
Figura 2 - “O Turco”	34
Figura 3 - Linha do Tempo.....	35
Figura 4 - Chaturanga	35
Figura 5 - Shatranj.....	36
Figura 6 - Neurônio e Neurônio Artificial.....	53
Figura 7 - Epigenética.....	66
Figura 8 - Lobos cerebrais	70
Figura 9 - Funções Executivas	72
Figura 10 - Teste de Turing com xadrez	79
Figura 11 - Identificação das coordenadas	87
Figura 12 - Identificação das peças	87
Figura 13 - identificação do movimento do Cavalo	88
Figura 14 - Possíveis movimentos para as Brancas no primeiro lance	89
Figura 15 - Possíveis movimentos para as Pretas no primeiro lance.....	89
Figura 16 - Opções de movimento para as Brancas	90
Figura 17 - Construção da árvore de decisão	92
Figura 18 - Aplicação da Função Avaliação no Ply 3	93
Figura 19 - Aplicação do algoritmo Minimax.....	93
Figura 20 – Alpha e Beta.....	94
Figura 21 – Poda Alpha-Beta	95
Figura 22 - Chunks	104
Figura 23 - Captura que foi um erro do computador.....	111
Figura 24 - A primeira escolha do Stockfish 18 não é a captura.....	112
Figura 25 - Composição de Penrose, 2017	113
Figura 26 - Rybka versus Nakamura, 2008	114
Figura 27 - Duchess versus Kaissa, 1977	115
Figura 28 - Rede Neural Convolucional.....	126
Figura 29 - Busca na Árvore de Monte Carlo	127
Figura 30 – Arquitetura do AlphaZero	128
Figura 31 - Fontes de Pesquisa	131
Figura 32 - Status dos artigos	131

Figura 33 - Nuvem de Palavras	132
Figura 34 - Diagrama 1 – escolha da peça do xeque-mate.....	143
Figura 35 - As cinco preferências do Stockfish no diagrama 1	144
Figura 36 - Diagrama 2 xeque-mate de Dama ou Torre.....	146
Figura 37 - Decisão do Stockfish 18 – Primeiro a Torre.....	147
Figura 38 - Diagrama 3 – Percepção Visual.....	149
Figura 39 - Decisão do Stockfish 18 – Cc3+ com empate	149
Figura 40 - Diagrama 4 – “Fair Play” de Magnus Carlsen.....	152
Figura 41 - Diagrama 5, Erro de Fischer	154
Figura 42 - Avaliação de Stockfish antes do lance de Fischer	155
Figura 43 - Avaliação de Stockfish depois do erro de Fischer.....	155
Figura 44 – Diagrama 6 - A Partida da Ópera	159
Figura 45 – Diagrama 7 - Posição de Penrose invertida.....	161
Figura 46 - Avaliação do Stockfish da posição de Penrose.....	162
Figura 47 - Diagrama 8 – Abstração e Analogia	166
Figura 48 - A posição correspondente ao Diagrama 8	166
Figura 49 - Memória de Trabalho, humano e computador.....	177

Lista de Quadros

Quadro 1 - Relações do Xadrez com as Funções Executivas.....	73
Quadro 2 - Abordagens da IA.....	75
Quadro 3 - Objeções de Turing	81
Quadro 4 - Atributo de valores materiais das peças de xadrez.....	97
Quadro 5 - Avaliação das Vantagens Estratégicas	97
Quadro 6 - Artigos encontrados.....	132
Quadro 7 - Relevância dos artigos.....	134
Quadro 8 - Artigos incluídos	135
Quadro 9 – Diagramas e funções cognitivas	140
Quadro 10 – O Erro de Fischer versus IA cometeria o mesmo erro?.....	158
Quadro 11 - Participante reconhece a partida versus força de jogo (rating)	160
Quadro 12 – Considerações a partir do JX sobre as objeções de Turing	173
Quadro 13 – Funções Cognitivas, <i>Engines</i> de Xadrez e Inteligência Humana	175

Lista de Gráficos

Gráfico 1 - Tempo que pratica o xadrez	140
Gráfico 2 - Participação em Torneios de Xadrez	141
Gráfico 3 - Leitura de Livros.....	141
Gráfico 4 - <i>Rating</i>	142
Gráfico 5 – Percentual da Escolha do Lance pelo Participante	144
Gráfico 6 - Distribuição de Lances por Faixa de Rating	145
Gráfico 7 - Percentual de Escolha do Participante	146
Gráfico 8 - Sacrifício de Dama ou Torre?	147
Gráfico 9 - Atenção Seletiva, percepção visual.....	150
Gráfico 10 – Entendeu o lance de Carlsen?.....	152
Gráfico 11 - Fischer cometeu um erro?	156
Gráfico 12 - A IA poderia cometer o mesmo erro de Fischer?	157
Gráfico 13 - Análise do “reconhecimento de padrão”.....	159
Gráfico 14 – Respostas ao diagrama 7	161
Gráfico 15 – Análise das Respostas do Diagrama 7.....	162
Gráfico 16 - Respostas do Diagrama 8.....	167
Gráfico 17 – Existe o xeque-mate no diagrama?.....	168

Lista de Abreviações

ADN – Ácido Desoxirribonucleico

CHC - Cattell, Horn e Carroll

CHOD - Conel Hugh O'Donel

ConvNet - Convolutional Neural Networks

CTS – Ciência, Tecnologia e Sociedade

DEC – Digital Equipment Corporation

DL – Deep Learning

DRL – Deep Reinforcement Learning

FIDE – Federação Internacional de Xadrez

GOF AI - Boa e Velha Inteligência Artificial

HLAI - Human-Level Artificial Intelligence

IA – Inteligência Artificial

IAF – Inteligência Artificial Forte

IAG – Inteligência Artificial Geral

IBM - International Business Machines

ITEP - Theoretical and Experimental Physics

JX – Jogo de Xadrez

LC0 – Leela Zero

MACHACK - Multiple Access Computing Machine Aided Cognition

MCTS – Monte Carlo Tree Search

MIT – Massachusetts Institute of Technology

ML – Machine Learning

QI – Quociente de Inteligência

RNA – Rede Neural Artificial

TAR – Teoria Ator-Rede

ToM – Teoria da Mente

WISC - Wechsler Intelligence Scale for Children

SUMÁRIO

1.INTRODUÇÃO.....	17
1.1 Apresentação	17
1.2 Problema.....	18
1.3 Hipótese.....	20
1.4 Questão Orientadora	20
1.5 Objetivos.....	20
1.5.1 Objetivo geral	20
1.5.2 Objetivos específicos.....	21
1.6 Justificativa.....	21
1.7 Metodologia.....	23
1.7.1 A Teoria Ator-Rede	23
1.7.2 Conceitos da Teoria Ator-Rede	24
1.7.3 Críticas à Teoria Ator-Rede.....	27
1.7.4 Aplicação da Teoria Ator-Rede.....	28
1.7.5 O Caminho da Rede Sociotécnica desta Tese	30
2.LINHA DO TEMPO: INTELIGÊNCIA ARTIFICIAL E XADREZ	34
3.A HISTÓRIA DA INTELIGÊNCIA ARTIFICIAL E XADREZ.....	52
4.A CONSTRUÇÃO SOCIAL DA INTELIGÊNCIA ARTIFICIAL	55
5.INTELIGÊNCIA E FUNÇÕES COGNITIVAS.....	65
5.1 Conceitos	65
5.2 Teorias	66
5.3 As Bases Biológicas da Inteligência.....	69
5.4 Funções Executivas e a Inteligência.....	71
5.5 Teoria da Mente, Atenção Seletiva, Memória Semântica, Memória Episódica e Emoções	72
6.MÁQUINAS BENÉFICAS	74
7.AS MÁQUINAS PODEM PENSAR?	79
8.COMO PENSAM OS COMPUTADORES DE XADREZ.....	84

8.1 Teoria dos Jogos	84
8.2 O Artigo de Claude Shannon.....	84
8.3 O Algoritmo Minimax.....	86
8.4 O Algoritmo Alpha-Beta Pruning	93
8.5 A Função Avaliação	96
8.6 Técnicas de Avaliação	97
9. COMO PENSAM OS MESTRES DE XADREZ.....	101
9.1 O Investigador Pioneiro, Adrian De Groot.....	101
9.2 Chunks.....	103
9.3 A importância de Conhecer o Pensamento do Mestre de Xadrez	105
9.4 Como Pensa Gary Kasparov.....	105
10. O LABORATÓRIO JOGO DE XADREZ.....	110
11. DEEP BLUE E KASPAROV	116
12. A REVOLUÇÃO DO ALPHAZERO	125
13. REVISÃO BIBLIOGRÁFICA	131
13.1 Resultados da Revisão	131
13.2 A Humanização do Algoritmo de Xadrez	135
13.3 Atribuição de Emoções aos Algoritmos de Xadrez.....	136
13.4 Análises de Algoritmos e de Humanos em Posições de Xadrez	136
13.5 Conceitos sobre o Xadrez e o Reconhecimento de Padrões.....	137
13.6 Avaliações Humanas são Relativas e de <i>Engines</i> Concretas.....	137
13.7 Discussões sobre a Revisão de Artigos	138
14. RESULTADOS DA PESQUISA SURVEY.....	140
14.1 Perfil dos Participantes	140
14.2 Diagrama 1, a Escolha Estética do Xeque-mate.....	142
14.3 Diagrama 2, o Sacrifício de Dama.....	145
14.4 Diagrama 3, Percepção Visual.....	148
14.5 Diagrama 4, “ <i>Fair Play</i> ”, o Contexto.....	151
14.6 Diagrama 5, Bobby Fischer Errou?	153

14.7 Diagrama 6, Reconhecimento de Padrão.....	158
14.8 Diagrama 7, Raciocínio Lógico e Abstração.....	160
14.9 Diagrama 8, Abstração	163
15. ENTREVISTAS COM MESTRES DE XADREZ, CIENTISTAS DA COMPUTAÇÃO E NEUROCIÊNCIAS	169
16. RESULTADOS e DISCUSSÃO	173
16.1 Análise da dimensão sociotécnica	173
16.2 Análise da dimensão cognitiva	175
16.2.1 Diferenças.....	175
16.2.2 Simetrias	176
16.2.3 Controvérsias	178
16.3 Análise da dimensão algorítmica.....	179
17. CONCLUSÕES.....	180
LIMITAÇÕES DESTA TESE	185
TRABALHOS FUTUROS	186
REFERÊNCIAS BIBLIOGRÁFICAS.....	187
APÊNDICES.....	196
APÊNDICE A – ENTREVISTA GM01	196
APÊNDICE B – ENTREVISTA GM02	204
APÊNDICE C – ENTREVISTA CC01	213
APÊNDICE D – ENTREVISTA CC02	223
APÊNDICE E – ENTREVISTA NC01	231
APÊNDICE F – ENTREVISTA NC02	242
APÊNDICE G – TCLE 1	251
APÊNDICE H – TCLE 2	254
APÊNDICE I – Perguntas CC	257
APÊNDICE J – Perguntas NC	258
APÊNDICE K – Perguntas MX	259
APÊNDICE L - SURVEY.....	260

APÊNDICE M – CARTA CONVITE.....	265
ANEXOS.....	267
ANEXO A	267
Parecer consubstanciado do CEP – Comitê de Ética em Pesquisa com Seres Humanos da UFSCar	267

1. INTRODUÇÃO

1.1 Apresentação

O termo inteligência artificial (IA) é por si só um tema controverso.

A princípio, IA é uma jornada da busca do ser humano pela sua essência. Este caminho desperta emoções, ninguém fica indiferente. É uma reflexão filosófica do nosso eu, de nossa busca pela superação. É a luta pela vida, pela existência.

Em nossos dias, a IA tem gerado preocupações aos humanos com a possível perda do trabalho, de ser excluído, de surgir uma sociedade de robôs, o trabalho automatizado, mecânico, desvalorizado; o capital privilegia apenas alguns.

A reflexão traz luz às diferenças, às culturas, à luta pela sobrevivência. O trabalho filosófico emerge da descoberta da própria essência do ser humano, de sua capacidade em construir ciência e tecnologia, sua criatividade, a superação de seus limites.

O astronauta do filme 2001, em sua volta para a espaçonave, após ter sido jogado no espaço, no universo sem fim, pela sua máquina inteligente, de ter encarado a morte, ter lutado e de ter resistido, de ter superado todas as adversidades, exerceu a sua engenhosidade, conseguiu retornar para a espaçonave e enfrentou a criatura, até desligar suas memórias e sua energia vital. Criador e criatura.

Nos capítulos que seguem, surge a engenhosidade humana em construir inteligência, construir processos. Conhecer a história, os laboratórios dos cientistas, os conflitos sociais da IA.

Por isso é necessário que esta leitura quebre paradigmas e se deixe levar pela busca incessante pela sua essência, pelo que de mais importante carrega dentro de si, não apenas sua inteligência, mas, principalmente, suas emoções, que a inteligência não é pura razão, mas também emoções. Não é apenas um cálculo frio, gélido. Não se trata apenas de memória, mas sim de uma combinação de memórias, de associações, de criatividade, de descoberta, da admiração – aquela lampadinha que se acende e que nos faz admirar o mundo que nos cerca - nossa própria existência e evolução, nosso encontro com nós mesmos, nossas esperanças, nossas angústias, medos. Convido os leitores para esta viagem de reflexão.

Quando eu tinha meus 16 para 17 anos, a professora de português pediu uma redação, escrevi com o título “O Homem na Lua”. Era sobre os gastos altíssimos com tecnologias enquanto a miséria imperava pelas ruas do planeta.

Depois veio nova redação, que intitulei “As Máquinas”. Eu comecei dizendo dos enormes benefícios dos aparelhos domésticos, como enceradeiras, máquinas de lavar; aspirador de pó; também dos carros, aviões, etc. E terminava a redação dizendo que, no confronto da época, entre o computador (Belle) e o campeão de xadrez Bobby Fischer, o computador perdeu feio. Pois o computador não tinha emoções.

Jamais podia imaginar que, tantos anos depois, estaria eu aqui escrevendo uma tese de doutorado, inspirado que fui pela minha orientadora, Sylvia, sobre um tema que caminhou ao meu lado durante todo este tempo e que estava adormecido. A vida deu voltas.

Percorrer os caminhos da IA, de suas origens, da inspiração com as neurociências é o que recomenda Bruno Latour, para a observação dos cientistas em seu trabalho, da entrada do laboratório e a descoberta das controvérsias, da ciência em ação, da construção social da tecnologia e dos valores encapsulados na interação humano não humano, homem e objeto, dentro de uma rede sociotécnica.

Como refletir sem cair no puro tecnicismo?

Alan Turing, no seu artigo “Podem as máquinas pensar?”, nos convida a uma reflexão, a filosofar, não a um puro esmiuçamento matemático da computação em busca de algoritmos cheios de cálculos e estatísticas do comportamento humano, da mente humana; ele também menciona nas entrelinhas a solidariedade que buscava para seus próprios anseios e angústias.

Licença poética e estilo, emitindo opinião, mostrando que a tese não trata de mero tecnicismo acadêmico, mas sim de uma reflexão sobre a própria essência do ser humano.

1.2 Problema

A inteligência artificial, desde seus primórdios, utilizou o jogo de xadrez como um de seus instrumentos nos laboratórios, buscando simular aspectos da inteligência humana. O xadrez tornou-se um ambiente privilegiado para investigar processos de tomada de decisão em condições controladas.

O xadrez é um ambiente determinístico, com informação perfeita, regras estáveis simples e claras. No xadrez os jogadores observam o mesmo estado do tabuleiro a cada lance da partida. Além disso, o espaço de estados possíveis é muito amplo, da ordem de 10^{43} , enquanto o número de partidas possíveis é da ordem de 10^{120} (número de Shannon)

o que o torna um desafio combinatório extremo, exponencial na profundidade da árvore, porém finito.

Assim, é possível estudar estratégias como: busca em espaço de estados, heurísticas, planejamento de múltiplos passos, e mais recente, técnicas de aprendizado de máquina e redes neurais.

Um marco histórico nesse percurso ocorreu em 1997, quando o sistema *Deep Blue*, desenvolvido pela IBM (*International Business Machines*), derrotou o então campeão mundial de xadrez Garry Kasparov. O evento foi amplamente interpretado como uma vitória simbólica da inteligência artificial sobre a inteligência humana em um dos domínios mais emblemáticos do raciocínio estratégico.

Entretanto, é fundamental observar que os processos utilizados por sistemas como o *Deep Blue* diferem fundamentalmente daqueles empregados por jogadores humanos. Enquanto o sistema computacional baseava-se predominantemente em força bruta combinada com heurísticas especializadas e avaliação massiva de posições por segundo, o jogador humano lida com reconhecimento de padrões, intuição, experiência acumulada, memória semântica, além de processos emocionais e regulatórios.

A vitória de sistemas computacionais sobre o campeão mundial de xadrez consolidou a narrativa da inteligência artificial superar os humanos em um de seus jogos mais antigos e complexos.

No entanto, permanece aberta a controvérsia se essa superioridade no jogo de xadrez implica em superioridade nos processos cognitivos humanos. Em outras palavras, o computador ter vencido o campeão mundial de xadrez significa que a inteligência artificial superou a mesma capacidade cognitiva dos humanos?

Apesar da vasta literatura sobre o desenvolvimento técnico das *engines* de xadrez e dos estudos cognitivos sobre expertise enxadrística, observa-se uma lacuna analítica quanto à comparação sistemática entre os processos de tomada de decisão humanos e dos algoritmos de xadrez no mesmo ambiente experimental. De um lado, a ciência da computação concentra-se predominantemente no desempenho e na eficiência dos algoritmos; de outro, a psicologia cognitiva e as neurociências investigam reconhecimento de padrões, memória e funções executivas em jogadores especialistas. Contudo, raramente esses campos são articulados em uma análise integrada que examine, de forma comparativa e interdisciplinar, como os agentes, humano e artificial, decidem, quais mecanismos mobiliza e quais implicações epistemológicas decorrem dessas

diferenças. Essa ausência de integração teórica e empírica justifica a necessidade do presente estudo.

1.3 Hipótese

A Inteligência Artificial e a Inteligência Humana, no laboratório do jogo de xadrez, não são equivalentes de um modo geral, nem possuem uma hierarquia epistemológica. Trata-se, pois, de sistemas cognitivos de naturezas distintas que exploram o mesmo espaço formal por meio de estratégias de tomada de decisão diferentes, configurando formas distintas de inteligência.

Parte-se da hipótese de que a superioridade das *engines* no desempenho enxadrístico não decorre da simulação dos processos cognitivos humanos, mas da adoção de uma racionalidade algorítmica específica, baseada em árvores de decisão, em otimização probabilística e exploração combinatória orientada por avaliação estatística. Por outro lado, a cognição humana trabalha por reconhecimento de padrões, heurísticas construídas pela experiência e aprendizagem, além de processos de regulação emocional e memória contextual. Nesse sentido, as *engines* de xadrez devem ser compreendidas como um artefato sociotécnico, cuja racionalidade emerge de condições históricas, técnicas e científicas.

1.4 Questão Orientadora

Quais são as diferenças, simetrias e controvérsias nos processos de tomada de decisão da inteligência artificial e dos seres humanos no jogo de xadrez?

1.5 Objetivos

1.5.1 Objetivo geral

Investigar as diferenças, simetrias e controvérsias nos processos de tomada de decisão entre a inteligência artificial e a inteligência humana no laboratório do jogo de xadrez.

1.5.2 Objetivos específicos

1. Analisar os processos de tomada de decisão da Inteligência Artificial no laboratório jogo de xadrez;
2. Examinar os processos de tomada de decisão da Inteligência Humana no jogo de xadrez;
3. Comparar as diferenças, simetrias e controvérsias entre a Inteligência Artificial e a Inteligência Humana nas tomadas de decisão no laboratório jogo de xadrez.

1.6 Justificativa

Apesar da vasta literatura sobre o desenvolvimento técnico das *engines* de xadrez e dos estudos cognitivos sobre expertise enxadrística, ainda se observa uma lacuna quanto à comparação sistemática e integrada entre os processos de tomada de decisão humanos e algorítmicos em um mesmo ambiente experimental.

De um lado, a ciência da computação concentra-se predominantemente no desempenho, na eficiência e na otimização dos algoritmos; do outro lado, a psicologia cognitiva e as neurociências investigam reconhecimento de padrões, memória, funções executivas e regulação emocional em jogadores de xadrez especialistas. Contudo, tais campos ainda estão distantes em análises comparativas que examinem os mecanismos decisórios de ambos os agentes e suas implicações epistemológicas.

A contribuição original desta tese consiste em integrar estes três níveis analíticos frequentemente tratados de maneira isolada, quais sejam: o funcionamento algorítmico das *engines* de xadrez; os processos cognitivos humanos envolvidos na tomada de decisão; a interpretação sociotécnica dessas diferenças no campo Ciência, Tecnologia e Sociedade (CTS).

Ao articular revisão sistemática da literatura, *survey* com análise de padrões decisórios e entrevistas com especialistas, o estudo propõe que o laboratório jogo de xadrez (JX) constitui um espaço epistemológico privilegiado para problematizar o próprio conceito de inteligência.

Embora as *engines* tenham superado os humanos no desempenho enxadrístico, cujo marco histórico foi o confronto entre *Deep Blue* e Garry Kasparov, com a vitória do computador, essa superioridade não decorreu da simulação dos processos cognitivos

humanos, mas da construção de uma racionalidade algorítmica específica, baseada em otimização probabilística e exploração combinatória.

As inteligências artificiais especialistas no JX são sistemas computacionais que realizam cálculos de acordo com árvores de decisões, estatísticas e probabilidades. Já os enxadristas tomam decisões ao realizar cálculos mentais de variantes de jogadas¹, avaliações provenientes do estudo, do aprendizado, de suas memórias de longo prazo, da regulação emocional e de suas experiências. Os mestres de xadrez não veem peças isoladas, eles veem estruturas significativas (*chunks*), o que reduz drasticamente o espaço de busca mental. Essas distinções suscitam controvérsias quanto à natureza da inteligência e à narrativa de substituição do humano pela tecnologia.

CTS é um campo interdisciplinar que investiga como ciência e tecnologia são produzidas e discutidas dentro dos contextos sociais, políticos, econômicos e culturais. Nesta tese, o campo CTS trata de conhecer a inteligência artificial dentro do contexto sociotécnico, em que a narrativa da superação da cognição humana no xadrez trouxe controvérsias sobre mudanças sociais, a substituição do ser humano pela tecnologia. Por outro lado, conhecer os processos de tomar decisões dos seres humanos no JX implica na inspiração de ideias de como a integração destas duas forças de inteligência podem ser direcionadas para solucionar outros aspectos da vida cotidiana, por exemplo, em procedimentos médicos, industriais, etc.

No Capítulo 11, dedicado ao confronto entre Garry Kasparov e o supercomputador *Deep Blue*, foram mobilizados os conceitos da Teoria Ator-Rede (TAR) para analisar as influências sociais, políticas, econômicas e tecnológicas envolvidas na construção e no desenvolvimento da inteligência artificial (IA). O embate representou um marco sociotécnico que impulsionou significativamente as pesquisas em IA, ao evidenciar que os avanços científicos não ocorrem de forma isolada, mas emergem de redes complexas de atores humanos e não humanos. As análises fundamentadas na TAR contribuem para o desenvolvimento do raciocínio crítico e para a compreensão da construção sociotécnica das revoluções científicas descritas por Thomas Kuhn, reforçando a relevância do campo CTS para a reflexão epistemológica e para a compreensão crítica da produção científica contemporânea.

¹ Na terminologia enxadrística, fazer cálculos durante a partida de xadrez significa imaginar as variantes que serão jogadas, os cálculos mentais de análise de variantes.

O objetivo que pode ser alcançado pelas ciências e tecnologia é de uma inteligência híbrida, ou de máquinas benéficas, em que aconteça a cooperação entre sistemas computacionais e cognição humana, em que um aprende com o outro, como foi o exemplo do xadrez e do jogo do Go.

1.7 Metodologia

1.7.1 A Teoria Ator-Rede

Para investigar o problema de pesquisa dentro do campo CTS, foi utilizada a metodologia da Teoria Ator-Rede (TAR), a partir das ideias de Bruno Latour, abordadas em seu livro “Cogitamus” (Latour, 2016).

A TAR é uma metodologia que auxilia na compreensão de ciência e tecnologia que formam parte de uma rede de atores, humanos e não humanos na construção social das tecnociências (Latour, 2011; 2012; 2013). Os laboratórios dos cientistas são examinados, o contexto social, o resultado de testes, as inscrições dos vários instrumentos e as consequentes controvérsias que surgem à medida em que vão sendo descritas e narradas (Latour, 2011; Latour e Woolgar, 1997). O pesquisador em CTS deve se manter imparcial e buscar os pontos de vista de forma simétrica e reflexiva (Latour, 2012; 2013), não importando nem os acertos nem os erros das ciências e tecnologias, pois “o que é verdadeiro relativamente ao ‘objeto’, o é ainda mais relativamente ao sujeito” (Latour, 2017, p. 228).

Michel Callon, Bruno Latour e John Law contribuíram para o desenvolvimento da TAR, a partir da publicação de seus artigos que investigaram as relações sociais das ciências e tecnologias (Sismondo, 2010; Bijker, 1995).

Callon tem dois trabalhos considerados clássicos utilizando a TAR: a construção do carro elétrico pelos engenheiros da Electricité de France com as células de combustível elétrico e que receberam muitas críticas sobre a viabilidade técnica pelos colegas da Renault, que trabalhavam no projeto de motores a combustão interna (Callon, 1987). O outro artigo de Callon (1986) foi sobre os pescadores, biólogos e as vieiras na baía de Saint Brieuc, França. A associação entre humanos (biólogos, pescadores) com não humanos (as vieiras) e suas respectivas perspectivas e interesses tiveram grande repercussão, com muitas críticas de outros pesquisadores.

Os trabalhos clássicos de Latour são a respeito da construção do motor à Diesel (Latour, 2011), e sobre os micróbios e Louis Pasteur (Latour, 2017). No relato sobre Diesel, Latour não deixa escapar de sua narrativa, atores como, engenheiros, cientistas, consumidores, financiadores e empresários, mas também o querosene e as bombas de propulsão. No relato sobre Pasteur, Latour encontra atores como militares, cientistas químicos e físicos, mas também vacinas e colônias de micróbios, além das inscrições geradas dentro do laboratório de Pasteur. As colaborações públicas recebidas por Pasteur garantiram ao cientista a adesão às suas ideias, pela relevância nas áreas da saúde, na ordem pública e no exército. Os militares franceses de 1880 estavam interessados em recrutar melhores soldados e Louis Pasteur traduziu esse interesse, via trabalho retórico, em apoio ao seu programa de pesquisa. Após o trabalho de Pasteur, os militares tiveram um novo interesse na pesquisa básica sobre micróbios.

Por sua vez, John Law (1987) descreveu em sua narrativa histórica as expansões portuguesas que só obtiveram sucesso porque os portugueses souberam associar os recursos políticos e financeiros com instrumentos astronômicos de navegação, cartas celestes, e em especial técnicas de navegação, recentes à época.

1.7.2 Conceitos da Teoria Ator-Rede

Dentre as técnicas utilizadas pela TAR, a etnometodologia é o instrumento para que se façam as agudas observações e descrições abordadas pelo pesquisador, e assim construir uma “cartografia das controvérsias” que constituirá a rede sociotécnica investigada. Nas observações dos pesquisadores não escapam detalhes da política, mitos, religiões, ritos e demais aspectos pouco explorados, pois não há lugar para preconceitos; o contexto social influi no que acontece nas pesquisas científicas (Latour, 2013).

A narrativa de uma rede sociotécnica é tecida com todos estes elementos e suas controvérsias, além dos atores humanos e não humanos que o pesquisador vai encontrar ao percorrer pelos fios condutores da rede em formato de labirinto. É justamente nas observações e descrições realizadas no percurso desta rede sociotécnica que surgem os bons relatos, segundo Latour, atento que está o pesquisador, primeiro com as controvérsias decorrentes das relações, ainda que momentâneas, entre atores; em segundo lugar, no distanciamento de espaço-tempo, que permite a inclusão de elementos exóticos (arcaicos, misteriosos); em terceiro lugar em distinguir os pontos de ruptura, a quebra de paradigma; e por último, a pesquisa em arquivos, documentos de jornais, revistas, museus

e outras mídias, além de entrevistas com os especialistas e outros atores que podem estar nos bastidores de conferências, quando surgem os prós e contras sobre os experimentos - as controvérsias (Latour, 2012; Latour, 2016). As metáforas na escrita e a retórica são os ingredientes que fornecem “sabor” e “colorido” que torna a narrativa do pesquisador um bom relato.

A linha condutora da narrativa são as controvérsias, o momento de ruptura, de revolução, da quebra de paradigma. Identificar estes momentos de possíveis fracassos ou acertos, e se colocar dos dois lados da disputa de forma simétrica, relativista, crítica e assim perceber o momento importante onde a “caixa preta” da tecnociência se fecha, findam as controvérsias e a ciência se estabelece (Latour, 2011).

Sobre o conceito de “caixa preta”, Latour (2000) conceitua como a tecnologia já estabelecida, que está pronta e que os aspectos de sua construção são aceitos e entendidos pela sociedade, portanto, de domínio público. É preciso pois, ir ao laboratório e investigar a ciência em ação, o trabalho desenvolvido pelos cientistas e técnicos, antes do fato consumado, antes que a caixa preta da ciência e da tecnologia se feche.

Os laboratórios de investigação onde os fatos científicos ocorrem são dos mais diversos, pois a TAR é uma metodologia adaptável, flexível aos locais onde se realizam os experimentos, sejam eles a baía de Saint Brieuc (Callon, 1986), o instituto Salk (Latour e Woolgar, 1997), a floresta amazônica (Latour, 2017) ou mesmo em um tabuleiro de xadrez, como no relato aqui apresentado mais à frente. Collins (1995), por exemplo, indica que os computadores são laboratórios naturais de estudos dos algoritmos. É suficiente, para serem laboratórios, que nestes locais existam os objetos que geram as provas científicas, as inscrições, que se tornarão a base para relatos e artigos (Latour, 2011).

As provas geradas nos laboratórios devem ser conhecidas e investigadas pois são a partir delas que os artigos são publicados, os debates se iniciam, as dúvidas acontecem e as controvérsias afloram. Controvérsias: “... todas as posições possíveis, que vão desde a dúvida absoluta ... até a certeza indiscutível” (Latour, 2016, p. 79). As provas nada mais são do que os fatos que acontecem nos laboratórios e analisar como elas (provas, fatos) são construídos é o que causou muitas polêmicas com os cientistas, que não admitiam a ideia de que fatores sociais poderiam influenciar os resultados de seus experimentos.

Os instrumentos dentro dos laboratórios são os não humanos que geram informações e sociabilizam com os cientistas as inscrições que formam as provas, que por

sua vez geram as dúvidas para outros pesquisadores, que por fim desencadeiam controvérsias (Latour, 2012).

Knorr-Cetina (1995) analisa as muitas influências que são exercidas sobre os cientistas dentro dos laboratórios, como por exemplo os recursos financeiros para suprir salários, materiais, equipamentos e outros custos, a organização política da instituição e suas relações com o Estado.

Com o uso da etnografia, os pesquisadores entram nos laboratórios e os examinam como locais de construção do processo de conhecimento. Assim fizeram Latour e Woolgar (1997) em seu livro “Vida de Laboratório”, onde investigaram não só o trabalho dos pesquisadores, documentos, arquivos, artigos, mas também as provas que eram geradas por instrumentos de pesquisa; questionaram, colocaram dúvidas, examinaram resultados.

Os resultados dos experimentos que ocorrem nos laboratórios são comunicados em artigos e a habilidade de escrita dos pesquisadores, seu poder de persuasão, as alianças com outros pesquisadores, a captação de recursos, são integrantes da prova científica, ou seja, os fatos científicos são construídos socialmente, já não são apenas “fatos naturais”, mas sim fatos moldados pela cultura (Knorr-Cetina, 1995).

Knorr-Cetina (1995) ainda cita uma lápide de um dos primeiros estudos de laboratório, atribuído a Doroty Sayers, na qual são comparados fatos (provas científicas) com vacas: se você olha fixamente para eles, geralmente fogem.

As ações entre atores da rede sociotécnica deslocam outros atores, existe um desvio do caminho previsto, consequência das suas intrincadas relações e tudo isto é observado e analisado na TAR, os rearranjos, as transformações sociais dos objetos, instituições, atores. Este conceito sobre o deslocamento – ou as translações, ou mesmo traduções, das relações entre humanos e não humanos, que em suas movimentações deslocam, transformam outros atores, todos eles conectados a uma rede, é um dos conceitos fundamentais da TAR, (Bjker, 1995; Latour, 2012). Na TAR, os objetos, as tecnologias, enfim, os não humanos trazem uma nova visão de como ciência e tecnologia são socialmente construídos.

Callon (1995) define o conceito de tradução como sendo as ligações entre dispositivos técnicos, teorias e seres humanos. A união destas traduções conduz para as redes de tradução, que variam em tamanho e complexidade e que ocorrem, sejam nos processos de associação, sejam na obtenção temporária das relações já estabilizadas.

Surge então um novo conceito neste modelo de tradução estendida, o ator passa a ser actante, ou, aquele que tem a capacidade de agir, as traduções que ocorrem, a história da qual estão participando naquele dado momento e contexto social e assim, a existência precede a essência. A tradução, no sentido da TAR não é neutra, os objetos carregam valores (Callon, 1995; Sismondo, 2010).

1.7.3 Críticas à Teoria Ator-Rede

Apesar de reconhecer a ampla possibilidade de utilização da TAR para narrar muitos casos, Sismondo (2010) critica o materialismo relacional do qual ela se utiliza, o que significa que suas aplicações são muitas vezes contraintuitivas. A narrativa sobre a construção social elaborada na TAR, o seu materialismo relacional, implica em dizer que os cientistas não fazem descobertas das propriedades naturais das coisas, mas sim, que as coisas e os fatos são resultado da construção social advindos de redes sociotécnicas.

Afirma Sismondo que a TAR vai contra as ideias intuitivas sobre as tecnologias terem propriedades reais e intrínsecas, não importam as relações sociais e em que rede social estejam. Mesmo que os interesses de um objeto possam ser manipulados, eles resistem a essa manipulação e, portanto, se opõem à rede. O realismo implícito da TAR passa a ser um contraponto com relação ao sucesso do relativismo metodológico (Sismondo, 2010).

Sismondo (2010) continua sua crítica ao considerar que a TAR se apoia apenas em escolhas racionais dos cientistas e engenheiros, deixando de lado os aspectos subjetivos de práticas e culturas. Ou seja, mesmo para escolhas racionais de engenheiros e cientistas, segundo Sismondo, os recursos são adequados para suas práticas dentro de seu contexto cultural a fim de atingir seus objetivos. No entanto, para a TAR, práticas e culturas precisam ser compreendidas em termos de arranjos dos atores que as produzem.

Sismondo (2010) critica também as figuras centrais das narrativas TAR, os protagonistas, e exemplifica dizendo que cientistas e engenheiros são mais parecidos com heróis ou heróis fracassados, ou mesmo pretensos heróis (e vilões?), perdendo assim atores secundários, citando as mulheres que foram marginalizadas nas ciências e tecnologias. E mais, os não humanos carecem de intencionalidade na maioria das vezes, sendo que, para ser um actante, estes devem agir, ter intenção. Se a TAR tem que negar que seja necessária a intencionalidade nas ações dos actantes e negar as diferenças entre humanos e não humanos em busca da simetria, que é um dos pilares para a teoria como

um todo, o princípio da simetria fica prejudicada, pois humanos e não humanos são tratados de modos diferentes. A riqueza de estratégias e interesses dos humanos são muito maiores do que a dos não humanos, sendo assim a TAR objetiva as ações dos cientistas e engenheiros e não é simétrica quanto aos não humanos (Sismondo, 2010).

Outra das críticas à TAR e seus atores-rede é pela dificuldade em explicar os acordos feitos pelos cientistas que resulta na ausência de controvérsias. Quando os cientistas se unem para entrar em acordo sobre os fatos e as interpretações alcançadas, o fazem para obter autoridade cognitiva e social. Nestas ocasiões, os conflitos desaparecem e os pesquisadores da TAR que observam os movimentos dos atores (no caso, os cientistas) acabam por não encontrar o elemento essencial, as controvérsias (Martin e Richards, 1995).

Outra crítica versa sobre a TAR ignorar e ocultar as ambivalências dos actantes, ou seja, os momentos em que eles surgem em outros pontos da rede sociotécnica em novas associações, em novas translações, deixando-os apenas com “super interpretações”. As mudanças e as novas associações entre atores são interpretadas como relativamente pequenas em um equilíbrio de elementos que entram em conflito com suas várias identidades, que estariam em redes sociais concorrentes (Wynne, 1995).

1.7.4 Aplicação da Teoria Ator-Rede

Em seu artigo sobre os métodos utilizados em CTS, John Law (2017) indica que os “estudos de caso” são formas instrutivas de modelo TAR para elaboração de narrativas das tecnociências. A utilização da TAR permite dar ênfase em se analisar as práticas de cientistas em seu contexto social, e que envolvem teorias, métodos e materiais, diferentes em diferentes locais, como por exemplo empresas, indústrias e laboratórios.

Dentre os exemplos de estudos de caso da TAR citado por Law, está a narrativa inovadora de Michel Callon, tendo como local de estudos a baía de Saint Brieu, que aborda a rede sociotécnica formada por cientistas, pescadores e vieiras, em vista da diminuição da população deste molusco, ou seja, a intrínseca relação entre atores humanos e não humanos, suas relações sociais, seus interesses, suas culturas, suas práticas. Os conceitos da TAR, portanto, estão em meio a narrativa.

A TAR é uma metodologia flexível, que se adapta aos locais dos experimentos, mesmo sendo eles em um JX ou em programas de computador. Collins (1995), por exemplo, indica que os computadores são laboratórios naturais de estudos dos algoritmos.

É suficiente, para serem laboratórios, que nestes locais existam os objetos que geram as provas científicas e as inscrições que se tornarão a base para relatos e artigos (Latour, 2011).

Como observou Collins (1995), há muito que se aprender com este novo laboratório chamado computador, onde são processados algoritmos especialistas, no caso softwares (programas de computador) que jogam xadrez. As novas tecnologias de IA que utilizam as chamadas redes neurais (*deep learning e machine learning*) que "aprendem" com os dados, são campos de estudos e testes de laboratório que tornam interessantes o valor de suas descobertas.

O Capítulo 11 é um exemplo da aplicação da TAR e seus conceitos enquanto acontecimento paradigmático para a compreensão das relações entre ciência, tecnologia e sociedade. Mais do que um episódio circunscrito à história da computação ou do xadrez, o embate constitui um marco sociotécnico no qual se articulam disputas simbólicas, interesses institucionais, processos de legitimação científica e reconfigurações epistemológicas acerca da inteligência, da cognição e da autonomia decisória das máquinas. Nesse contexto, a utilização TAR, desenvolvida sobretudo por Bruno Latour, Michel Callon e John Law, permite compreender o episódio não como resultado linear do progresso técnico, mas como efeito de uma rede heterogênea de mediações envolvendo atores humanos e não humanos, tais como cientistas, engenheiros, corporações, algoritmos, infraestrutura computacional, mídia especializada e imaginários sociais sobre a inteligência artificial.

A partir dessa perspectiva, o confronto Kasparov versus *Deep Blue* ultrapassa a dimensão competitiva do jogo de xadrez e assume relevância epistemológica no interior do campo CTS, ao evidenciar que os processos de produção científica e inovação tecnológica são indissociáveis das dinâmicas sociais, econômicas, políticas e culturais que os constituem. O avanço das pesquisas em inteligência artificial após o evento demonstra que determinadas rupturas tecnocientíficas operam como catalisadoras de novos regimes de visibilidade, financiamento, credibilidade e reorganização paradigmática do conhecimento científico.

Tal interpretação dialoga diretamente com as formulações de Thomas Kuhn acerca das revoluções científicas, especialmente no que se refere às transformações paradigmáticas que alteram os critérios de validação, os modelos explicativos e os horizontes epistemológicos de uma comunidade científica. Entretanto, ao incorporar os

pressupostos da TAR, esta análise desloca a centralidade exclusivamente cognitiva dos paradigmas kuhnianos para enfatizar a dimensão sociotécnica das controvérsias científicas, evidenciando que a estabilização do conhecimento decorre de processos de negociação, tradução e associação entre múltiplos atores.

Desse modo, o capítulo 11 busca contribuir para uma compreensão crítica da IA não apenas como artefato tecnológico, mas como construção sociotécnica inscrita em redes complexas de poder, conhecimento e mediação, reafirmando a relevância do campo Ciência, Tecnologia e Sociedade para a investigação contemporânea das transformações científicas e tecnológicas.

1.7.5 O Caminho da Rede Sociotécnica desta Tese

O caminho percorrido nesta rede sociotécnica é apresentado na Figura 1. Neste caminho são feitas pesquisas em repositórios de artigos científicos, livros, audiovisuais. Entrevistas com especialistas, das áreas da ciência de computação, das neurociências e com mestres de xadrez; *survey* (levantamento de dados) junto à comunidade dos amadores de xadrez e experimentos com *engines* concluem a tarefa da investigação para se obterem conclusões a respeito da pergunta de pesquisa. No capítulo 15 são apresentadas as diferenças entre IA e IH, as simetrias, controvérsias, os prós e contras.

Três métodos foram utilizados durante a tese: uma revisão sistemática da literatura sobre o tema das tomadas de decisão das inteligências artificiais e o JX; entrevistas semiestruturadas com mestres de xadrez, cientistas da computação e neurocientistas; um *survey* com amadores de xadrez para que fossem obtidas observações etnográficas.

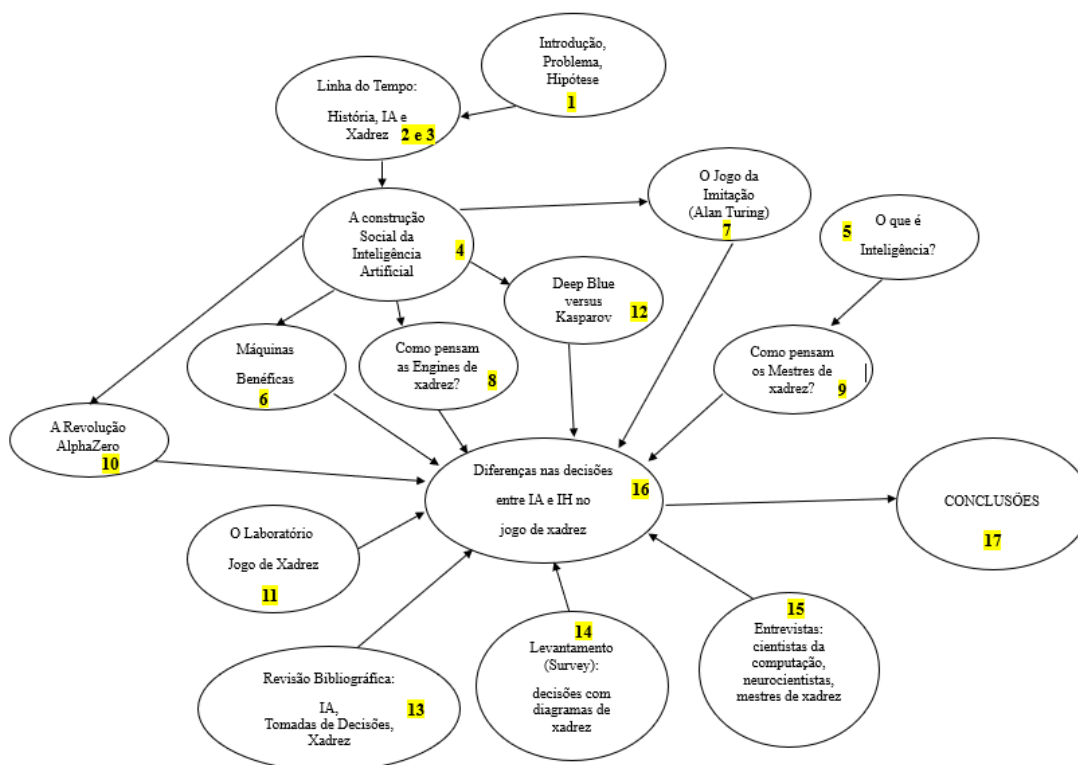
A revisão sistemática (Gil, 2002) permite a contextualização do problema, uma base com a fundamentação teórica, além do conhecimento de investigações anteriores. Os artigos encontrados pela revisão desta tese contribuíram com ideias para a elaboração dos diagramas do *survey*.

A entrevista semiestruturada (Gil, 2002) é a técnica de questionar o pesquisado em uma relação de pontos de interesse que são explorados pelo entrevistador de forma flexível; este usa de estratégias para se obter dados específicos com perguntas previamente elaboradas para a condução da entrevista. As entrevistas realizadas com os mestres de xadrez oportunizaram questionar como é o processo de tomada de decisão nas partidas e outras nuances sobre as *engines*. No caso dos cientistas da computação foi

abordado temas como os de abstrações e analogias e o futuro das pesquisas de IA com o xadrez. Com os neurocientistas foi abordado aspectos cognitivos que influem no processo de tomada de decisões e assuntos relacionados com as contribuições em áreas de pesquisa entre IA e neurociências.

Survey (Babbie, 2001) é um método de coleta de dados que serve para ampliar o conhecimento sobre um grupo social, no caso, uma amostra dos enxadristas amadores de diferentes lugares. Com esta amostra foi possível fazer observações etnográficas sobre, por exemplo, a força de jogo dos praticantes de xadrez e o reconhecimento de padrões.

Figura 1 - Os passos metodológicos da pesquisa



Fonte: o autor, 2026.

O percurso na rede sociotécnica inicia com a linha do tempo com a cronologia da história da IA, de algoritmos e computadores relacionados com o JX.

Metodologia:

- 1) Introdução: apresentação, problema, hipótese, questão orientadora, objetivos e justificativa;
- 2) Linha do Tempo. Um panorama cronológico de pontos marcantes da história da IA e suas ligações com o xadrez, situando historicamente acontecimentos importantes tanto do xadrez quanto da IA e em especial da integração dos dois.

- 3) História da IA com xadrez. Conhecer de forma breve a história e as relações entre a IA e um de seus principais laboratórios, o JX.
- 4) Relato sobre a construção social da IA. Alinhar a história da IA com as influências sociais, em especial quando de suas origens.
- 5) Pesquisa bibliográfica sobre inteligência; definições e conceitos, suas bases biológicas e sociais que servem de fundamentação para se conhecer as diferenças, simetrias e controvérsias entre a IH e a IA.
- 6) Máquinas benéficas e o xadrez; ou como a IA e o JX podem ser utilizados fora de competições, em contextos em que a IA pode apoiar os humanos em aspectos sociais de saúde e educação.
- 7) Artigo de Alan Turing, “as máquinas podem pensar?”; as objeções de Turing servem de base para examinar as diferenças entre IA e humanos aplicadas ao JX. A partir das objeções de Turing e da troca da proposição, “as máquinas podem pensar” para o “jogo da imitação”, é possível relacionar a ideia de que a IA é uma construção sociotécnica.
- 8) Conhecer como o computador pode jogar xadrez; o artigo seminal de Claude Shannon e os algoritmos Minimax e a poda Alpha-Beta. A função avaliação é apresentada e a engenhosidade dos humanos em construir heurísticas e implementá-las nas *engines* de xadrez.
- 9) Conhecer os elementos sobre o raciocínio humano para o JX, o pioneirismo de Adrian De Groot e sua pesquisa com reconhecimento de padrões e um resumo de como pensa o campeão de xadrez Kasparov para tomar decisões durante uma partida;
- 10) As redes neurais e o AlphaZero, o algoritmo que revolucionou o xadrez e o jogo do Go.
- 11) Alguns exemplos de jogos de programas de xadrez, tanto contra humanos quanto contra outros programas que trazem alguns indicadores sobre as diferenças entre IA e IH.
- 12) A narrativa do confronto entre *Deep Blue* e Kasparov, considerado o marco histórico da IA e seus aspectos sociotécnicos. Aqui é apresentado como as *engines* de xadrez foram uma construção sociotécnica que envolveu o marketing da IBM, truques psicológicos dos programadores de *Deep Blue* para

confundir Kasparov, erros de computação e toda a pressão social que envolveu o confronto.

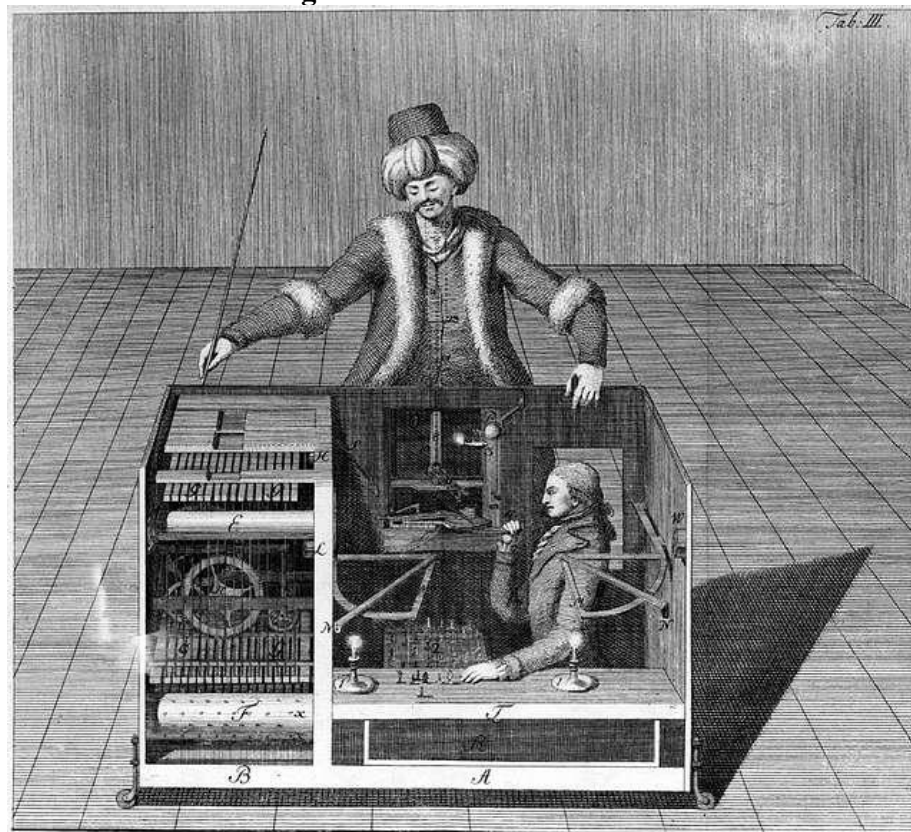
- 13) Revisão sistemática envolvendo o tema da IA, as tomadas de decisão e o JX, onde se busca conhecer o estado da arte referente a pergunta de pesquisa.
- 14) Apresentação dos resultados da aplicação do *survey* com diagramas de xadrez para entender processos de tomada de decisão de jogadores de xadrez amadores em temas específicos e aplicá-los a apreciação da *engine* Stockfish.
- 15) Entrevistas com os mestres de xadrez, cientistas da computação e neurocientistas abordando os temas desta tese e suas controvérsias.
- 16) Resultados, análises e discussões da pesquisa, apresentando observações sobre as funções cognitivas, comparações entre *engines* e humanos, as diferenças, simetrias e controvérsias e a elaboração argumentativa a partir do artigo de Alan Turing da construção sociotécnica de *engines* de xadrez.
- 17) Conclusões, contendo as revelações e confirmações sobre os experimentos obtidos, correlacionando a força dos jogadores de xadrez com o reconhecimento de padrões, do processo cognitivo atencional; a associação entre memória de trabalho e memória de longo prazo de *engines* e humanos; a construção sociotécnica da IA; a ideia do autoconhecimento do ser humano ao comparar a sua inteligência com a IA.

2. LINHA DO TEMPO: INTELIGÊNCIA ARTIFICIAL E XADREZ

Preâmbulo

Em plena sociedade dos anos 1770, para surpresa da população europeia, surge um “autômato” (Figura 2), uma máquina que continha uma mesa com o tabuleiro de xadrez e um boneco, a semelhança de um “turco”, com seus olhos expressivos e sua cabeça protegida por um turbante. Qual tecnologia seria possível para que uma simples máquina, abarrotada de engrenagens no interior de sua caixa, fosse capaz de desafiar o ser humano em um jogo tão complexo como o xadrez? O computador só iria surgir 150 anos depois (Figura 3).

Figura 2 - “O Turco”



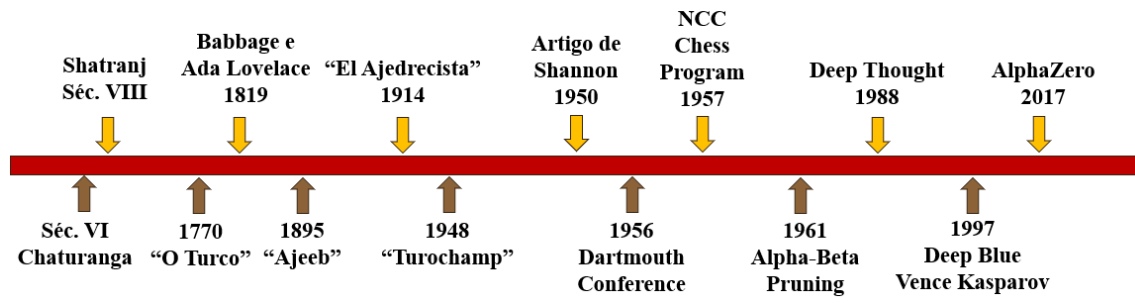
Fonte: Creative Commons:

<https://picryl.com/media/the-turk-engraving-in-joseph-friedrich-racknitz-ueber-den-schachspieler-des-2299a7>

Cronologia

Figura 3 - Linha do Tempo

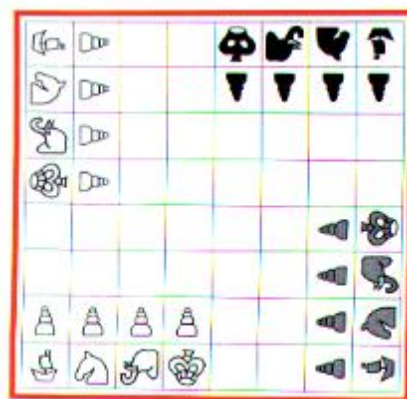
Linha do Tempo: Inteligência Artificial e Xadrez



Fonte: o autor, 2026

Século VI: Surge o xadrez. O JX é um dos jogos mais antigos da humanidade e foi desenvolvido ao longo dos séculos por diversos povos (Murray, 1913). O primeiro antecessor do JX, o jogo indiano chamado "Chaturanga" (Figura 4), era jogado por quatro adversários representando os quatro elementos: fogo, terra, água e ar e para cada lance da partida eram jogados dados. Esta associação com os elementos da natureza e a casualidade da sorte condizem com aspectos mais ligados a origens místicas do jogo (Murray, 1913).

Figura 4 - Chaturanga

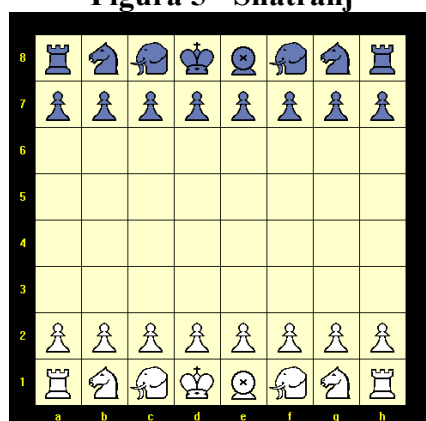


Fonte: Elaboração gráfica da Universidade Federal do Paraná (UFPR) a partir de modelo conceitual do autor, (dez. 2020).

Século VII e VIII: Uma das revoluções do JX aconteceu com os persas e depois com os árabes, que modificaram algumas regras importantes, dentre elas, a eliminação dos dados que estavam presentes até aquele momento, tornando o jogo puramente de estratégia. A outra mudança foi com relação ao número de jogadores, reduzidos a dois, juntando-se os exércitos de peças dois a dois, que passaram a se enfrentar em uma luta direta (Murray,

1913). Esta mudança ocorrida no jogo, originando o “Shatranj” (o nome dado pelos árabes) é uma mudança de paradigma, um desvio dos valores do “Chaturanga”. O “Shatranj” (Figura 5) era um jogo onde apenas a estratégia do mais hábil prevalecia. O pensamento lógico passou a ser associado com o JX, e os árabes o utilizaram como instrumento de educação da matemática (Shenk, 2007).

Figura 5 - Shatranj



Fonte: Elaboração gráfica da Universidade Federal do Paraná (UFPR) a partir de modelo conceitual do autor, (dez. 2020).

Século IX – XVI: Ocorreram na Europa várias pequenas revoluções sucessivas com o “Shatranj”, transformando-o no JX medieval. As peças mudaram suas formas antigas para a representação europeia da época. O Shah se tornou o Rei, surgiram as Torres, Bispos. Por último, a Rainha (Dama), espelhando a força que poderosas governantes, como Isabel de Castela, Catarina de Médicis, Catarina de Aragão, Maria Tudor, Elizabeth I da Inglaterra, Maria da Escócia, possuíam em suas épocas. O tabuleiro ganhou os contrastes de casas claras e escuras; as regras de movimentação das peças se modificaram para dar mais dinâmica ao jogo, refletindo os tempos renascentistas. O novo desenho das peças de xadrez foi uma metáfora da força política simbólica adequada aos povos europeus (Shenk, 2007).

Século XIII: É publicado o livro do rei Alfonso X com regras e problemas do JX (Murray, 1913; Shenk, 2007).

Século XVI: Livro de xadrez de Ruy Lopes que vai unificar as regras do xadrez (Murray, 1913).

Século XVII: A semente embrionária das ideias de uma tecnologia para realizar cálculos se iniciam com o ábaco, instrumento com mais de 3.000 anos de idade. Porém, a automatização através de uma máquina de cálculos só ocorreu no século XVII, com

Wilhelm Schickard em 1623 e com Blaise Pascal e sua máquina de calcular, a “Pascalina”, em 1642². Posteriormente, a evolução destas máquinas de calcular passaram por grandes nomes da matemática, como o de Gottfried Leibniz e Charles Babbage.

1770: O inventor-relojoeiro, o alemão Barão Von Kempelen, desenvolveu uma máquina incomum, que jogava xadrez com a habilidade de um mestre. Esta máquina, batizada de “O Turco”, foi a sensação das cortes europeias no final do século XVIII e início do século XIX. O “Turco” era requisitado em vários países e jogou contra Napoleão Bonaparte. Em um episódio, o imperador francês tentou trapacear no jogo contra o “Turco”, que teria tido um “rompante de cólera” ao derrubar as peças do tabuleiro diante de Napoleão (Wairy, 1895 apud Müller e Schaeffer, 2018). Características humanizadas de computadores se iniciaram ali. O “Turco” era uma caixa cheia de engrenagens com um boneco em cima que movia as peças. Boneco este que escondia habilidosos jogadores de xadrez, estes sim, que imprimiam inteligência e movimento às jogadas do “Turco”, que iludiram e deslumbraram os nada céticos espectadores da época (Shenk, 2007).

1819: Charles Babbage conhece o autômato jogador de xadrez, “O Turco” em Londres, e retorna no ano seguinte para jogar uma partida com a máquina. Anos depois desenvolve sua máquina analítica de calcular, o precursor dos computadores. Babbage imaginava que sua máquina poderia fazer cálculos para jogar xadrez. Sua máquina recebia “inputs” e gerava “outputs” - cartões perfurados -, armazenava em memória e recebia instruções. Enquanto Babbage pensava na arquitetura de sua máquina de calcular, Ada, Condessa de Lovelace, pensava em como instruir a máquina para executar tarefas. Ada Lovelace é considerada a primeira programadora da história da computação, criadora dos “algoritmos”. Babbage pensou em testar sua máquina com algoritmos que jogassem xadrez. Babbage pode ser considerado o pai da computação; Ada Lovelace, a mãe dos programas de computador (Müller e Schaeffer, 2018).

1850: Torneio de Londres, que define as atuais regras do xadrez e se inicia a utilização do relógio em todas as competições (Kasparov, 2004; Shenk, 2007).

1895: O autômato “AJEEB” (sucessor do “Turco”) disputa uma partida contra o campeão do torneio de Hastings de 1895, Harry Pillsbury (Charles Hooper, Royal

² <https://www.computerhistory.org/revolution/calculators/1/47>

Pollytechnical Institute, 1868, (apud Müller e Schaeffer, 2018)). AJEEB terminou seus dias em 1929, em um incêndio, tal qual aconteceu com seu antecessor, “O Turco”, nos EUA em 1854 (Müller e Schaeffer, 2018).

1899: Alfred Binet, francês, intrigado em descobrir como funcionava a inteligência dos enxadristas, realizou experimentos no final do século XIX com mestres de xadrez (Shenk, 2007). Binet acreditava na “memória fotográfica” dos enxadristas. Após seus experimentos, concluiu que eram armazenados nas suas memórias, modelos abstratos e não uma “fotografia” do tabuleiro (Nicolas, 1994; Shenk, 2007).

1912: Ernst Zermelo publica um artigo que origina o algoritmo Minimax (Russell e Norvig, 2024).

1914: É apresentada a máquina "El Ajedrecista" que executava a tarefa de xeque-mate com Rei e Torre contra Rei, formulada por Leonardo Torres y Quevedo, inspirado nas ideias de Charles Babbage. “El Ajedrecista”, construído em 1912, pode ser considerada a primeira máquina que jogava xadrez (Müller e Schaeffer, 2018).

1930 ... 1950: Uma grande revolução emerge com novas tecnologias, em especial com os grandes computadores. As máquinas de calcular podem, desde então, ser programadas com códigos de instruções para executar tarefas de propósito geral, bastando que fossem programadas para tal. As tarefas podem ser “comportamentos”, linha metodológica da escola “behaviorista” do início do século XX, cujos grandes cientistas eram Skinner, Pavlov, Throndicke (Shenk, 2007; Müller e Schaeffer, 2018).

1938 ... 1945: A Segunda Guerra Mundial; Alan Turing inicia os trabalhos em sua máquina "Enigma", junto a outros matemáticos, entre eles o campeão inglês de xadrez, Conel Hugh O'Donel Alexander (CHOD), para quebrar o código secreto dos alemães (Müller e Schaeffer, 2018).

1940: Claude Shannon trabalha nos laboratórios da Bell; Shannon é considerado o pai dos "sistemas da informação" (Müller e Schaeffer, 2018).

1941: Konrad Zuse reivindica a construção do primeiro computador programável o “Z3”; idealiza como desenvolver um algoritmo para que seu computador jogue xadrez (Müller e Schaeffer, 2018).

1946: Em suas pesquisas, Adrian De Groot revelou que a memória para o JX “reconhece padrões”, confirmada por Connors *et al.* (2011) que replicaram seu experimento (Gobet, 2019).

1946 – 1950: Alan Turing e John Von Neumann dois grandes gênios matemáticos, desbravadores da computação, deram um passo importante para o que conhecemos hoje como computadores (Russell e Norvig, 2024).

1948: Alan Turing escreve um algoritmo para jogar xadrez, a partir de ideias que teve em suas discussões enquanto nos trabalhos da guerra no Bletchley Park, com seus colegas CHOD Alexander, Harry Golombek e Stuart Milner-Barry (Müller e Schaeffer, 2018; Turing, 1988);

Norbet Wiener, matemático do Massachusetts Institute of Technology (MIT): “poderia ser construída uma máquina de xadrez que poderia muito bem ser um jogador tão bom quanto a grande maioria da raça humana” (Wiener, 1948, apud Müller e Schaeffer, 2018, p. 34, tradução nossa³).

Fred Reinfeld (enxadrista) em seus comentários de uma partida realizada em Varsóvia, 1935, entre Saviely Tartakower e Lajos Steiner:

Quando o professor Wiener do MIT inventou uma máquina de calcular que requer somente um décimo milésimo de segundo para as mais complicadas computações, ele disse, entre aspas ‘eu desafio você a descrever uma capacidade do cérebro humano que não posso duplicar com dispositivos eletrônicos’. Até o momento em que estas linhas foram escritas, o professor aparentemente ainda não havia aperfeiçoado um dispositivo eletrônico capaz de realizar movimentos de xadrez como o 20º lance de Tartakower... No entanto, este dia ainda pode chegar, quando veremos livros como ‘Os 1.000 Melhores Jogos de Robô’, ou quando torneios de xadrez terão que ser adiados devido à escassez de aço. (Reinfeld, 1948, apud Müller e Schaeffer, 2018, p. 35, tradução nossa⁴).

1949: Claude Shannon faz uma palestra onde anuncia como uma máquina poderia jogar xadrez (Müller e Schaeffer, 2018).

1950: Claude Shannon elabora os princípios de como deve ser um algoritmo de xadrez

³ “A chess machine could be built that “might very well be as good a player as the vast majority of the human race.”.

⁴ “When Professor Wiener of the Massachusetts Institute of Technology invented a calculating machine which requires only one ten-thousandth of a second for the most complicated computations, he was quoted as saying, ‘I defy you to describe a capacity of the human brain which I cannot duplicate with electronic devices.’ “Up to the time these lines were written, the Professor had apparently not yet perfected an electronic device capable of making such chess moves as Tartakower’s 20th... The day may yet come, however, when we shall see such books as Robot’s 1000 Best Games, or when chess tournaments will have to be postponed because of a steel shortage”.

em seu artigo "Programing a computer to play chess" (Shannon, 1988; Muller e Schaeffer, 2018).

Alan Turing publica o artigo "Computing Machinery and Intelligence", com o conhecido "Teste de Turing", ou o "Jogo da Imitação". Foi o passo decisivo para o que hoje é chamada "Inteligência Artificial". Turing é considerado o pai da IA (Turing, 1950; Russell e Norvig, 2024);

1951: Dietrich Prinz escreve um algoritmo para resolver problemas de xadrez de xeque-mate em dois lances no estilo "força bruta". Trata-se da primeira validação de resolução de problemas de mate em 2 por computadores (Muller e Schaeffer, 2018).

Marvin Minsky e Dean Edmonds desenvolveram a primeira rede neural artificial nomeada de SNARC, utilizando 3.000 válvulas para simular uma rede contendo 40 "neurônios" (Russell e Norvig, 2024).

1952: O jogo de Damas tem seu primeiro algoritmo escrito e que funcionava, de Christopher Strachey, o primeiro programa de computador de um jogo (Müller e Schaeffer, 2018).

Alan Turing testa seu algoritmo em uma partida de xadrez, apenas "no papel", inspirado pelas ideias de Claude Shannon (Müller e Schaeffer, 2018).

1953: Alan Turing anota a segunda partida disputada com a simulação de seu algoritmo (Müller e Schaeffer, 2018);

1954: Morre Alan Turing, aos 41 anos de idade.

1956: Dartmouth AI Conference: John McCarthy, Marvin Minsky, Nathaniel Rochester, Ray Solomonoff e Claude Shannon (também presentes Herbert Simon, Allen Newell e Arthur Samuel). A primeira menção ao termo "Inteligência Artificial". Um dos desafios iniciais que foi discutido e debatido no evento foi o de construir um programa de computador que jogasse bem xadrez (McCarthy *et al.* 1956).

O nascimento da era do software com a linguagem Fortran (Müller e Schaeffer, 2018).

No laboratório de pesquisas de Los Alamos, James Kister, Paul Stein, Stanislaw Ulan, William Walden, Mark Wells, tem a ideia de construir um programa de computador para

jogar xadrez no computador MANIAC I, o Los Alamos Chess Program (Müller e Schaeffer, 2018).

Apesar das limitações de regras e tabuleiro, pela primeira vez na história um programa de xadrez venceu um humano (Müller e Schaeffer, 2018).

A IBM publica anúncio buscando pessoas interessadas em se desenvolver como programadores de computador, cujo fundo do anúncio é a figura de um cavalo preto de xadrez com os dizeres: "para todos aqueles que gostam de jogar xadrez e resolver quebra-cabeças irão gostar deste trabalho" (Müller e Schaeffer, 2018).

1957: NSS Chess Program (Newell, Shall, Simon, 1988) programa de xadrez baseado em heurística.

Herbet Simon previu que em 10 anos um programa de xadrez seria campeão mundial (Müller e Schaeffer, 2018).

1958: Uma das primeiras partidas jogadas entre humano, Alex Bernstein, e um computador (IBM 704) (Horizons of Science, Vol. 1, nro 4; Education Testing Service, 1958; <https://www.youtube.com/watch?v=NUjiUR0ZH58&t=261s>).

O NSS chess program (Newell, Shall, Simon, 1988) derrota um humano no xadrez (Müller e Schaeffer, 2018).

Frank Rosenblatt desenvolveu o programa chamado “Perceptron” uma rede neural artificial (RNA) que podia aprender com os dados, sendo o precursor das atuais DL (Mitchell, 2019).

John McCarthy desenvolveu a linguagem de programação LISP (Müller e Schaeffer, 2018).

1959: Arthur Samuel em seu artigo cunha o termo “machine learning” (Mitchell, 2019; Russell e Norvig, 2024).

Oliver Selfridge publicou o artigo "Pandemonium: A Paradigm for Learning", contribuição histórica sobre ML, descrevendo um modelo para encontrar padrões em eventos (Russell e Norvig, 2024)

1961: O algoritmo Alpha-Beta de Alan Kotok (aluno de John MacCarthy) é uma melhora do algoritmo Minimax (Müller e Schaeffer, 2018).

1962: Alan Kotok, apresenta sua tese do algoritmo "Alpha-Beta" (Müller e Schaeffer, 2018).

O programa de computador que joga Damas, de Arthur Samuel vence uma partida contra o futuro campeão, do estado de Connecticut Robert Nealey. Foi a primeira vez na história de uma vitória de um computador contra um forte jogador (Müller e Schaeffer, 2018).

1965: Jonh McCarthy e o diretor do Institute for Theoretical and Experimental Physics (ITEP) de Moscou, Alexander Kronrod anunciam uma disputa de xadrez de 4 jogos entre os softwares americanos e russos (Müller e Schaeffer, 2018).

Alexander Kronrod (matemático russo): “O xadrez é a drosophila da inteligência artificial” (Ensmenger, 2012).

1966: O software para Damas de Arthur Samuel joga 4 partidas contra o então campeão mundial, Walter Hellman e contra o desafiante ao título, Derek Oldbury. O software perde todas as partidas (oito ao todo). Foi a primeira disputa em jogos de Damas entre um campeão mundial e um software.

O programa soviético de xadrez da equipe de Alexander Kronrod venceu com duas vitórias e dois empates, o programa norte-americano de xadrez da equipe de John McCarthy. O programa soviético se utilizava do tipo de programação “1” de Shannon, ou seja, a “força bruta”, o que indicava uma forte associação entre o poder computacional e a resolução do xadrez. No final das contas, Kronrod foi rebaixado no ITEP devido a reclamações dos seus superiores soviéticos de que muito tempo valioso de computação estava sendo desperdiçado com o xadrez (Müller e Schaeffer, 2018).

Joseph Weizenbaum criou o programa “Eliza”, que se envolvia em conversas com humanos (Russell e Norvig, 2024; Mitchell, 2019).

1967: Richard Gleenblat, do MIT, desenvolve seu próprio programa de xadrez, com base nos trabalhos de Alan Kotok, o MackHack VI (Projeto do MIT, o MAC – *Multiple Access Computing* para somas; *Machine Aided Cognition*, para pesquisadores em IA; além da junção de seu programa Hack, uma palavra em moda no MIT no início dos anos 1960).

O programa rodava em um *Digital Equipment Corporation* (DEC) PDP 6 (Gleenblat *et al.*, 1967, apud Müller e Schaeffer, 2018). O diferencial do MackHack é que ele possuía uma base de dados de aberturas e a segunda ideia foi uma tabela de transposição, ou seja, para posições iguais que eram calculadas em vários nós da árvore de decisão, o que evitava novos cálculos desnecessários. Além disso, a busca na árvore de decisão tinha uma rotina que verificava as chamadas “posições silenciosas” (onde ocorriam sequências de trocas de peças). No torneio, MackHack perdeu as quatro primeiras partidas e na última, em um final de jogo, os programadores propuseram empate ao adversário que, temendo a força do computador nos finais de jogo, aceitou. Pouco depois o programa recebe um “*rating*”⁵ da federação dos EUA de 1.239 pontos, o que era suficiente para vencer a classe ‘D’ do torneio, portanto, MackHack VI recebeu um troféu de campeão da classe “D” do torneio. Estabeleceu-se o início da competição homem versus máquina (Müller e Schaeffer, 2018).

Professor Hubert Dreyfuss, do MIT, um dos maiores críticos da IA, perde uma partida de xadrez para o programa MacHack VI (Müller e Schaeffer, 2018).

Em seu segundo torneio, o campeonato do Estado de Massachussetts, depois de algumas correções no programa, MacHack VI vence a primeira partida contra um humano, em apenas 21 lances, realizando um xeque-mate, com sacrifício de Dama (Müller e Schaeffer, 2018).

1968: John McCarthy disse a David Levy que, em 10 anos, haveria um programa que venceria mestres de xadrez (Müller e Schaeffer, 2018; Kasparov, 2017).

1969: Marvin Minsky e Seymour Papert publicaram o livro *Perceptrons*, que descreveu as limitações de redes neurais simples e fez com que a pesquisa de redes neurais declinasse e a pesquisa de IA simbólica prosperasse (Mitchell, 2019).

1970: O Ex-campeão mundial de xadrez, Max Euwe, doutor em engenharia e matemáticas, declara a importância dos exemplos de programação com o JX como contribuição para o entendimento de como a mente humana funciona (Müller e Schaeffer, 2018).

⁵ *Rating* é a pontuação internacional dos jogadores de xadrez calculados pela FIDE; também conhecido pelo nome de seu desenvolvedor, Arpad Elo, ou *rating* ELO.

1971: John McCarthy recebe o prêmio Turing por suas contribuições no desenvolvimento da IA e suas ideias dos programas de xadrez (Müller e Schaeffer, 2018).

1972: Inovações soviéticas importantes nos algoritmos de xadrez: processar enquanto é a vez do adversário; bit-boards; busca com um movimento "nulo"; e busca por movimentos com analogias, ou de raciocínio "estilo humano" (Müller e Schaeffer, 2018).

1973: Baseados em De Groot, Simon e Chase introduziram um novo conceito sobre a memória, chamado de fragmentação (*chunks*). Simon e Chase compararam o xadrez com as ciências biológicas declarando ser o jogo a “*drosophila*” das ciências cognitivas (Simon e Chase, 1973), em alusão às pesquisas biológicas que utilizavam a mosca como paradigma de experimentos biológicos. Os pesquisadores fizeram uma melhoria do que havia feito De Groot em 1946: embaralharam as peças no tabuleiro, gerando posições aleatórias que deveriam ser memorizadas pelas diversas classes de jogadores. Neste teste engenhoso, as diferenças foram poucas entre os mais fortes jogadores e os mais fracos (Simon e Chase, 1973).

James Lighthill divulgou o relatório "*Artificial Intelligence: A General Survey*", que fez com que o governo britânico reduzisse significativamente o apoio à pesquisa de IA (Mitchell, 2019).

1975: Newell e Simon ganham o prêmio Turing devido às suas pesquisas em IA, em parte com suas pesquisas com o xadrez (Müller e Schaeffer, 2018).

1976: Primeiro computador comercial para jogar xadrez: MicroChess (Müller e Schaeffer, 2018).

1977: Devido a computadores mais velozes, os algoritmos de xadrez, no caso o Chess 4.5, conseguem vencer o primeiro jogador humano com mais de 2.000 pontos de *rating*. O Grande mestre Michael Stean perde para o algoritmo Chess 4.6 - a primeira derrota de um grande mestre de xadrez para um algoritmo (Müller e Schaeffer, 2018).

1978: O microcomputador, Sargon, venceu um computador de grande porte no xadrez (Müller e Schaeffer, 2018).

O Algoritmo Chess 4.6 passa a barreira dos 2.000 pontos de *rating* da Federação Internacional de Xadrez (FIDE, 2023) (Müller e Schaeffer, 2018).

1979: O algoritmo BKG9.8 de gamão vence o campeão mundial em uma melhor de 10 partidas. Foi a primeira derrota humana de um campeão mundial para um computador em um jogo clássico (Müller e Schaeffer, 2018).

Neil Charness (1992) quebrou o velho paradigma do jogador de xadrez prodígio; os mestres são formados com muito trabalho, não nascem prontos, eles aprendem a ser bons jogadores (Shenk, 2007).

1980: Acontece a primeira vitória de um algoritmo (MOOR) no jogo Othelo (Reversi) de um campeão mundial humano (Müller e Schaeffer, 2018).

Máquinas Lisp simbólicas foram comercializadas, sinalizando um renascimento da IA. Anos depois, o mercado de máquinas Lisp entrou em colapso (Müller e Schaeffer, 2018).

1982: O valor estimado mundialmente em vendas de computadores de xadrez era de 100 milhões de dólares. O então campeão mundial de xadrez, Anatoly Karpov venceu o programa Fidelity em uma simultânea (Müller e Schaeffer, 2018).

A velocidade de processamento é fundamental para a força dos computadores, aumentando 100 pontos de *rating* a cada vez que dobra a velocidade (Ken Thompson, 1982, apud Müller e Schaeffer, 2018). Cada meio lance (ply) de profundidade equivale a 250 pontos de *rating*.

1984: Marvin Minsky e Roger Schank cunharam o termo “Inverno da IA” em uma reunião da Associação para o Avanço da IA, alertando a comunidade empresarial que o entusiasmo pela IA levaria à decepção e ao colapso da indústria, o que aconteceu três anos depois (Mitchell, 2019).

1985: O computador “Hitech” faz os primeiros "reconhecimentos de padrões" (Müller e Schaeffer, 2018).

Judea Pearl introduziu a análise causal de redes bayesianas, que fornece técnicas estatísticas para representar incertezas em computadores (Russell e Norvig, 2024).

1986: Os algoritmos da retropropagação chamaram a atenção com o artigo seminal de David Rumelhart, Geoffrey Hinton e Ronald Williams, "Aprendendo representações por meio da retropropagação de erros" (Russell e Norvig, 2024; Mitchell, 2019).

1988: Gary Kasparov, então campeão mundial de xadrez e considerado um dos melhores de todos os tempos, comentou que mesmo nos anos 2000 (12 anos após esta afirmação) os computadores não teriam chances contra os fortes jogadores e ainda assim, ele estaria ali para fazer aconselhamentos. Bent Larsen, um dos melhores grandes mestres dos anos 60 e 70 perde para *Deep Thought* (Müller e Schaeffer, 2018).

Deep Thought, precursor de *Deep Blue*, faz sua estreia contra grandes mestres de xadrez com algumas vitórias (Müller e Schaeffer, 2018).

1989: Kasparov afirma: “Ridículo! Uma máquina sempre será uma máquina, uma ferramenta para ajudar o trabalho de preparação dos jogadores. Nunca serei derrotado por uma máquina. Nunca será inventado um programa que supere a inteligência humana. E quando digo inteligência, também quero dizer intuição e imaginação. Você consegue imaginar uma máquina escrevendo um romance ou uma poesia?” – em entrevista a Thierry Paunin, (*Jeux & Stratégie*, n. 55, p. 4-5, 1989 apud Müller e Schaeffer, 2018, p. 239, tradução nossa⁶).

Peter Brown convence os diretores a trazerem o projeto “*Deep Thought*” para a IBM, que vê uma oportunidade de atrair pessoas talentosas para as pesquisas e investir massivamente na mídia. A IBM fez uma oferta irresistível para a equipe de “*Deep Thought*”: ter acesso aos imensos recursos da empresa com o objetivo de o projeto de xadrez ter uma conclusão de sucesso. A IBM notou o enorme potencial comercial de um campeonato mundial homem versus máquina. Prenúncio do projeto “*Deep Blue*” (Müller e Schaeffer, 2018).

Gary Kasparov jogou (e venceu) duas partidas contra o computador “*Deep Thought*” (Müller e Schaeffer, 2018; Kasparov, 2017).

A primeira aplicação prática da retropropagação com Yann LeCun que mostrou a combinação de retropropagação com as “redes neurais convolucionais”, com o reconhecimento de dígitos manuscritos (Russell e Norvig, 2024; Mitchell, 2019).

⁶ Ridiculous! A machine will always remain a machine, that is to say a tool to help the player work and prepare. Never shall I be beaten by a machine! Never will a program be invented which surpasses human intelligence. And when I say intelligence, I also mean intuition and imagination. Can you see a machine writing a novel or poetry?

1990: Anatoly Karpov enfrenta *Deep Thought* e sai vencedor (Müller e Schaeffer, 2018).

(Hsu *et al.*, 1990, apud Kasparov, 2007) - O computador não imita o pensamento humano, porém, por diferentes vias, consegue excelentes resultados com o xadrez.

1992: O programa "Chinnok", (jogo de Damas), perde a primeira disputa pelo campeonato mundial Homem Máquina contra o então campeão mundial de Damas, Marion Tinsley (Müller e Schaeffer, 2018).

1993: A estreia de *Deep Blue*. A derrota contra o grande mestre dinamarquês Bent Larsen mostrou um dos “bugs” do programa. Larsen trocou seus bispos pelos cavalos de *Deep Blue* e fechou a posição, onde os programas não conseguem analisar posições estratégicas de largo alcance (Müller e Schaeffer, 2018).

Estudos com pessoas com muitos anos de treino em determinada área, como o xadrez e a música, e a correlação com a sua inteligência é conhecido como “prática deliberada”. A “prática deliberada”, portanto, poderia influenciar apenas parte da inteligência que são relacionadas com os estudos intensos, mas não com a parte criativa, o que, talvez, possa ser chamado de “talento”. Uma controvérsia de intensa disputa nas ciências cognitivas, há mais de um século (Gobet, 2019).

1994: Mikhail Botvinnik:

O programa de nosso laboratório é o único que não usa a “força bruta”. Nosso programa pensa de uma maneira similar como os mestres de xadrez. “*Deep Thought*” calcula 50 milhões de posições em 3 minutos. No entanto, o nosso programa busca apenas entre vinte ou trinta posições, da mesma maneira que os mestres de xadrez.”. “O oponente de um computador pode aprender a jogar xadrez enquanto joga contra o computador”. “No futuro, o computador será capaz de fazer análises. Quando isto acontecer, ninguém mais vai publicar análises sem consultar o computador. Isto vai mudar de forma drástica a literatura do xadrez (Müller e Schaeffer, 2018, p. 272-273, tradução nossa⁷).

⁷ This is the only program in the world that doesn't use brute force. Instead of using brute force our program “thinks” in a similar manner as a chess master thinks. DEEP THOUGHT analyses one hundred and fifty million positions in three minutes. They are working on a program that will look at two billion positions in three minutes. However, my program looks only at twenty or thirty positions, just as a chess master would do. the opponent of the computer can learn to play chess while playing the computer. Improved. In the future and that the computer will be able to make analyses. When that happens no one will publish analyses anymore without consulting the computer. This will drastically change the chess literature.

O programa Fritz derrota Kasparov em partidas "*blitz*" (de 5 minutos). Meses depois, Kasparov perde para outro programa, o Chess Genius, em uma partida rápida (25 minutos por jogador) (Müller e Schaeffer, 2018).

No jogo de Damas, o programa Chinnok empata seis partidas com o campeão do mundo Tinsley, que morreria um ano depois. Após Tinsley, nenhum humano derrotou Chinnok no jogo de Damas, que se aposentou em 1997 (Müller e Schaeffer, 2018).

1995: Com patrocínio da Intel, Kasparov joga duas partidas contra o programa Fritz, vencendo a primeira, depois de se constatar um erro do operador para o lance de Fritz (um erro humano!) e empatando a segunda partida (Müller e Schaeffer, 2018).

Joel Benjamin, grande mestre de xadrez dos EUA é convidado a participar do time de *Deep Blue* (Müller e Schaeffer, 2018).

1996: David Bronstein:

Hoje em dia, os computadores são tão inteligentes que conseguem fazer jogadas brilhantes que representam um desafio até para os grandes mestres! Eles não têm expectativas, nem alegria, nem decepção. Nunca se cansam. Não fazem ideia de quem são seus adversários e nem sequer temem alguém como eu, que um dia lutou pela coroa. (David Bronstein, apud Müller e Schaeffer, 2018, p. 268, tradução nossa⁸).

Acontece o esperado encontro entre o campeão mundial de xadrez Gary Kasparov e a *engine* da IBM, *Deep Blue*. *Deep Blue* possuía 256 chips capazes de analisar 150 milhões de posições no tempo de 3 minutos. *Deep Blue* vence a primeira partida, e fez história para a IA. Porém, Kasparov vence três partidas das cinco restantes, vencendo o *match* por 4 x 2 (3 vitórias, dois empates, uma derrota) (Müller e Schaeffer, 2018; Kasparov, 2007; Kasparov, 2017).

Fernand Gobet e Herbert Simon, citam dois processos cognitivos para se jogar xadrez: o reconhecimento de posições (memória de trabalho) e o planejamento (Gobet e Simon, 1996).

⁸ Now computers are so clever they can make brilliant moves that pose problems even to grandmasters! They have no expectations, no joy, no disappointment. They never tire. They have no idea who they are playing against and are not even afraid of someone like myself who once fought for the crown.

1997: A equipe de *Deep Blue* investiga os erros do encontro de 1996 entre a *engine* e Kasparov e percebem que é necessário programar o "entendimento" das posições de xadrez para poder vencer Kasparov (Müller e Schaeffer, 2018);

Quem sofreu as consequências da derrota para um computador foi nada mais nada menos que o então campeão mundial de xadrez e considerado um dos maiores jogadores de todos os tempos, Gary Kasparov. Seu adversário foi o supercomputador da IBM, “*Deep Blue*”. Sem dúvida, um marco histórico. A partir daquele momento a IA retomou com força o seu desenvolvimento, com as técnicas desenvolvidas de Aprendizagem de Máquina e RNA. (Müller e Schaeffer, 2018; Kasparov, 2007; Kasparov, 2017).

No documentário de Frank Marshall (2014) sobre o confronto, um dos componentes da equipe de *Deep Blue*, Joel Benjamin, responde sobre algumas questões que ficaram obscuras em 1997, como a de um lance quando o algoritmo entrou em *loop*.

1998: O programa de computador TDGamon de Gerauld Tesauro perdeu para o então campeão mundial do jogo de gamão, Malcon Davis, numa melhor de 100 partidas. Na revisão das partidas foi encontrado um “bug” no programa. Poucos anos mais tarde, o programa era superior ao ser humano (Müller e Schaeffer, 2018);

O programa MAVEN de Brian Sheppard derrota o campeão dos Estados Unidos, Adam Logan, no jogo de palavras-cruzadas (*scrabble*) (Müller e Schaeffer, 2018).

2002: David Bronstein afirma a necessidade de se jogar contra computadores, porque são os melhores professores (Müller e Schaeffer, 2018).

O programa de computador de Henri Bal e John Romein resolveu o jogo africano da Mancala (Awari), concluindo que as melhores jogadas resultariam em empate (Müller e Schaeffer, 2018).

2003: Kasparov enfrenta – e empata – em uma série de 6 partidas contra o programa “*Deep Junior*”. Posteriormente, Kasparov enfrenta o programa Fritz X3d e empata, desta vez em 2 a 2 (Müller e Schaeffer, 2018).

O grupo de pesquisa da Universidade de Alberta enfrenta e perde por uma estreita margem de pontos, um profissional de Poker Texas Hold’en, entre dois jogadores (Müller e Schaeffer, 2018).

2006: O IBM Watson surgiu com o objetivo inicial de vencer um humano no icônico programa de perguntas e respostas “Jeopardy!” (Müller e Schaeffer, 2018).

2007: Jonathan Schaeffer anuncia ter resolvido o jogo de damas de 64 casas, sendo que, com um jogo perfeito de ambos os lados, a partida será empate (Müller e Schaeffer, 2018).

2010: O ex-campeão mundial Vladimir Kramnik considera os programas de xadrez insuperáveis para os humanos (Müller e Schaeffer, 2018).

2011: O software Watson, da IBM, vence o programa de televisão “Jeopardy!”; o sistema de computador de perguntas e respostas derrotou o campeão (humano) de todos os tempos do programa, Ken Jennings (Müller e Schaeffer, 2018; Mitchell, 2019).

2012: Gary Kasparov vence o algoritmo inventado por Alan Turing, o “Turochamp” (Kasparov e Friedel, 2018).

2013: A DeepMind introduziu o aprendizado de reforço profundo, uma ConvNet que aprendia com base em recompensas e aprendia a jogar por meio da repetição, superando os níveis de especialistas humanos (Mitchell, 2019).

2015: O grupo de pesquisa da Universidade de Alberta anunciou que seu programa de computador resolveu o jogo do poker, Texas Hold’em para duas pessoas (Müller e Schaeffer, 2018).

2016: O AlphaGo da DeepMind derrotou o melhor jogador de Go, o campeão mundial Lee Sedol em Seul, Coreia do Sul, gerando comparações com a partida de xadrez de Kasparov com o *Deep Blue* quase 20 anos antes (Müller e Schaeffer, 2018; Mitchell, 2019; Silver *et al.*, 2018).

2017: O computador aprendeu a “jogar xadrez do zero”, apenas realizando partidas contra si mesmo. AlphaZero utiliza o sistema “Monte Carlo”, aliado a RNA e a “*deep reinforcement learning*” (Müller e Schaeffer, 2018; Mitchell, 2019; Silver *et al.*, 2018).

O físico britânico Stephen Hawking alertou: "A menos que aprendamos a nos preparar e evitar os riscos potenciais, a IA pode ser o pior evento na história da nossa civilização" (Mitchell, 2019, p. 15, tradução nossa⁹).

⁹ “The development of full artificial intelligence could spell the end of the human race”.

2018: Yann LeCunn, Geoffrey Hinton e Yoshua Bengio, recebem o Prêmio Turing, por suas contribuições fundamentais para o aprendizado profundo¹⁰

2022: O engenheiro de software do Google Blake Lemoine foi demitido por revelar segredos do LaMDA e alegar que ele possuía consciência (CNN Brasil, 2022).

2024: Geoffrey Hinton e John Hopfield recebem o prêmio Nobel de Física, por suas descobertas que permitem o ML com redes neurais artificiais¹¹.

Demis Hassabis e John Jumper, ambos da DeepMind recebem o prêmio Nobel de Química por seus trabalhos com o AlphaFold, para previsão da estrutura de proteínas¹².

¹⁰ <https://awards.acm.org/about/2018-turing>

¹¹ <https://www.nobelprize.org/prizes/physics/2024/summary/>

¹² <https://www.nobelprize.org/prizes/chemistry/2024/summary/>

3. A HISTÓRIA DA INTELIGÊNCIA ARTIFICIAL E XADREZ

As origens da IA têm como marco o encontro de pesquisadores em Dartmouth, 1956, quase todos matemáticos, envolvidos com a nova ciência que se iniciava. Logo neste primeiro encontro, em que John McCarthy citou pela primeira vez o termo “inteligência artificial” (McCarthy *et al.*, 1956), os pesquisadores já discutiam sobre o melhor método de implementar IA a um computador. E é neste ponto que se inicia a primeira controvérsia, afinal, o que é a inteligência e como ela poderia equipar as máquinas? Os principais nomes daquele encontro tentaram definir o que era inteligência, mas mesmo outras áreas do conhecimento, como a neurociência, a psicologia e a filosofia, ainda hoje se desafiam para definir o termo (Mitchell, 2019).

Sem esta primeira definição consolidada sobre o que é inteligência, de forma elaborada e explicada, vários métodos para se construir algoritmos que pretendiam simular a IH deram início a uma disputa entre grupos de pesquisadores por, principalmente, verbas governamentais ou de empresas interessadas nos resultados desta nova tecnologia revolucionária (Mitchell, 2019).

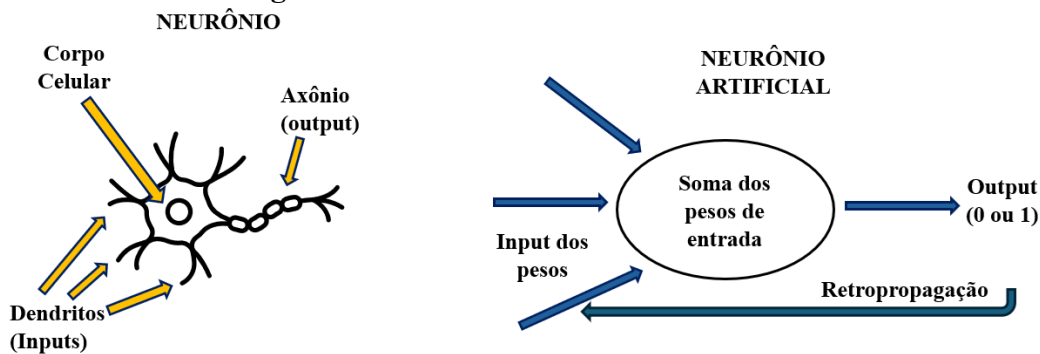
Dentre as ideias sugeridas, houve aqueles que propuseram que a IA deveria ser baseada no pensamento lógico, racional, dedutivo. Outros propuseram que o método deveria ser indutivo, baseado em estatísticas e probabilidades. Um terceiro método propunha que a IA deveria estar de acordo com a forma do cérebro humano funcionar, imitar as comunicações das células neurais, em uma inspiração vinda da biologia, que mostrou que o neurônio, muito semelhante ao que faz um computador, recebe informações (entradas ou *inputs*), processa e propaga um resultado para outros neurônios (saídas ou *outputs*) (Mitchell, 2019).

Nos anos seguintes ao encontro de Dartmouth, surgiu o algoritmo chamado “*Perceptrons*”, baseado no funcionamento da comunicação entre os neurônios¹³ e que está em evidência neste fim de quarto do século XXI, com as DL. Seu autor, o psicólogo Frank Rosenblatt, pensou em etapas no processamento de dados onde se atribuem pesos e valores que modificam os resultados a serem obtidos, atualizando-os por repetição, até se conseguir chegar em um resultado ideal. Este processo seria o aprendizado das redes neurais artificiais (RNA), inspiradas também pelas ideias da psicologia comportamental,

¹³ É importante deixar claro que o que é chamado de RNA é uma construção computacional que se compara ao sistema nervoso biológico apenas por uma analogia muito livre. De fato, entidades radicalmente distintas.

em especial às ideias do psicólogo behaviorista Burrhus Skinner, que entendia que comportamentos podem ser aprendidos e condicionados pela repetição e reforço positivo. Posteriormente, foi implementado os algoritmos de “retropropagação” (Figura 6), cujos autores foram David Rumelhart, Geoffrey Hinton e Ronald Williams em 1986 e a primeira aplicação da ideia veio com Yan LeCun em 1989 (Mitchell, 2019).

Figura 6 - Neurônio e Neurônio Artificial



Fonte: o autor, 2026

A ciência da computação precisava testar seus algoritmos. Os pioneiros da IA, dentre os quais Alan Turing, John McCarthy, Marvin Minsky, Herbert Simon, Alan Newell e Claude Shannon utilizaram o JX, por causa de suas regras simples e informação completa (Shannon, 1988; Turing, 1988). Todas as informações para se tomar uma decisão estão à mostra e não existe o elemento do acaso para invalidar o caminho racional, lógico, matemático.

O JX passou a ser um laboratório constante, não só da IA, mas também de outra ciência que emergiu com a computação, a ciência cognitiva. Dentre estes cientistas, são publicadas pesquisas de Herbert Simon e Allen Newell (Newell, 1988; Newell, Shaw e Simon, 1988). Em uma alusão aos experimentos com a mosca de laboratório das ciências biológicas, Simon e Chase (1973) cunharam o termo “drosophila das ciências cognitivas”, ou seja, o JX é o modelo experimental de laboratório onde podem ser estudados os processos cognitivos humanos, como memória, percepção e resolução de problemas (Charness, 1992), e que também foi utilizado pelos cientistas da computação para desenvolver a IA (Ensmenger, 2012). As redes neurais complexas, sejam as da IH, sejam as da IA, usufruíram das muitas possibilidades empíricas que o JX possibilitou aos cientistas.

Para o computador se tornar “inteligente”, bastaria que superasse os seres humanos em seu jogo de origem milenar, o xadrez. Assim que um algoritmo vencesse o

campeão mundial no JX, a IA se estabeleceria, seria o primeiro triunfo no caminho para se gerar inteligência a uma máquina.

Um marco para as pesquisas de IA aconteceu em 1997, com a vitória do computador “*Deep Blue*” da IBM, sobre o então campeão mundial de xadrez Gary Kasparov. *Deep Blue* foi uma arquitetura computacional aliada a um algoritmo especializado apenas em jogar xadrez, ou seja, uma *engine* especialista em xadrez (Ensmenger, 2012; Bory, 2019; Mitchell, 2019).

Sem dúvida, a vitória de *Deep Blue* foi um divisor de águas na história da IA.

Cèvora (2019) pontua que, enquanto *Deep Blue* conseguia jogar xadrez no nível técnico de Kasparov, por outro lado, não executava outras atividades concomitantes, como por exemplo, identificar e movimentar as peças no tabuleiro, que exige uma inteligência visual e motora tão importantes quanto o ato de conceber lances no xadrez.

A inteligência é mais do que resolver problemas isolados, é todo um planejamento de existência (Lauro, 2017). *Deep Blue* apenas jogava xadrez e o que se reconhece como inteligência está muito distante de apenas calcular variantes do jogo (Mitchell, 2019).

O objetivo dos pioneiros da IA era que esta pudesse superar os humanos em inteligência e a este conceito se denomina nos dias de hoje como inteligência artificial geral (IAG), ou, de outra forma, inteligência artificial forte (IAF) (Searle, 1984; Mitchell, 2019).

4. A CONSTRUÇÃO SOCIAL DA INTELIGÊNCIA ARTIFICIAL

A IA é construída a partir de algoritmos desenvolvidos pela IH. Esses algoritmos encapsulam lógicas formalizadas em linguagens computacionais e dão origem a ferramentas que auxiliam diversas atividades do cotidiano. Entre elas, destacam-se os programas para jogar xadrez, historicamente centrais para a pesquisa em IA (Russell e Norvig, 2024).

Como argumenta Ensmenger (2012), o xadrez de computador pode ser compreendido como a própria história da IA. O jogo funcionou como um laboratório privilegiado para o estudo de algoritmos de busca, avaliação e tomada de decisão. Nesse contexto, Collins (2010) observa:

O xadrez é um dos principais pontos de discussão sobre a noção de inteligência artificial. Hubert Dreyfus afirmou, de forma memorável, que nenhum computador jamais venceria um grande mestre de xadrez, pois era impossível definir e programar as regras utilizadas por esses mestres. Ele estava errado, pois um programa executado em um computador que ficou conhecido como *Deep Blue*, desenvolvido pela IBM, de fato venceu grandes mestres, e hoje sabemos que um computador suficientemente poderoso e programado adequadamente provavelmente vencerá o melhor jogador de xadrez humano na maioria das vezes. Mas há outro sentido em que Dreyfus não estava errado. Assim como no caso do ciclismo, é preciso observar atentamente a maneira como o ser humano, por um lado, e o computador, por outro, executam a tarefa. (Collins, 2010, p. 106; tradução nossa¹⁴).

Um jogo como o xadrez é um grande desafio pela sua complexidade, tanto nos cálculos de variantes do jogo, quanto na conceitualização dos padrões recorrentes que, muito fácil, são reconhecidos pelos mestres de xadrez para solucionar problemas (Collins, 2010).

A força bruta dos cálculos de variantes por um supercomputador e a indecifrável maneira como os mestres de xadrez reconhecem padrões de forma inconsciente são opostos entre computadores e humanos (Collins, 2010).

Para que os computadores pudessem fazer suas buscas exaustivas na árvore de decisão, deveriam fazer podas significativas para que fosse possível andar pelos ramos da

¹⁴ *Chess is one of the key points of discussion for the notion of artificial intelligence. Hubert Dreyfus famously claimed that no computer would ever beat a chess grand master because it was impossible to set out and program the rules used by grand masters. He was wrong in that a program running on a computer that came to be known as Deep Blue, developed by IBM, did beat grand masters, and we now understand that a powerful enough computer appropriately programmed is likely to beat the best human Chess player most of the time. But there is another sense in which Dreyfus was not wrong. As in the case of bicycle riding, one must attend carefully to the way the human on the one hand, and the computer on the other, carry out the task.*

árvore, apenas, um ou dois passos a mais que a capacidade que têm os seres humanos, mesmo sem conseguir chegar ao objetivo do jogo, o xeque-mate (Collins, 2010).

Os conceitos relacionados ao jogo, a possibilidade de avaliações heurísticas sofisticadas eram elementos importantes, mas depois soube-se que apenas poder chegar uns dois lances à frente dos humanos era suficiente para poder vencer uma grande parcela das disputas. Os cientistas da computação conseguiram seu intento por conta da vitória de *Deep Blue* sobre Kasparov, mas a maneira como a IH decidia, ficou longe dos sofisticados algoritmos de xadrez. “Se definirmos um bom jogo de xadrez simplesmente como ‘a capacidade de vencer’, então jogar xadrez é outro exemplo de conhecimento tácito de limite somático” afirmou Collins (2010, p. 108, tradução nossa¹⁵).

Há muita diferença entre a maneira como os humanos fazem as coisas com a natureza do próprio conhecimento. Criar programas que pudessem simular o cérebro humano, era uma meta por trás do que se pretendia com a IA. “No caso do xadrez, depende se ‘jogar xadrez’ significa ‘vencer no xadrez’ ou ‘vencer no xadrez usando reconhecimento de padrões do tipo humano’” (Collins, 2010, p. 119, tradução nossa¹⁶) Collins (2010) continua:

No xadrez de campeonato mundial, há muito mais em jogo do que apenas o que acontece no tabuleiro — há toda a estratégia que começa com a definição dos desafios e a escolha do local do torneio — na qual Bobby Fischer parece ter sido um grande expoente. O xadrez jogado nas circunstâncias da Guerra Fria não segue as mesmas convenções de estética e estratégia do xadrez da primeira metade do século XX, e provavelmente mudou novamente com a vitória do *Deep Blue*. Não existem máquinas capazes de gerenciar tais compensações e reparos ou aplicar as regras da estratégia de uma maneira humana e sensível ao contexto social. Nem é possível imaginar como inventar tais máquinas. Para entender como essas coisas devem ser feitas, precisamos nos envolver com a vida social. (Collins, 2010, p. 120, tradução nossa¹⁷).

Collins (2010) prossegue seu argumento enfatizando a necessidade do ser humano na produção do conhecimento: sem o agente humano, não há como socializar a máquina.

¹⁵ *If we define good chess playing simply as “the ability to win,” then Chess playing is another example of somatic-limit tacit knowledge.*

¹⁶ *In the case of chess it depends on whether “playing chess” was taken to mean “winning at chess” or “winning at chess by using human-type pattern recognition.*

¹⁷ *In world championship chess there is much more going on than what happens on the board—there is all the gamesmanship that begins with the setting up of the challenges and the location of the tournament—at which Bobby Fischer seems to have been an archexponent. Chess played under the circumstances of the cold war does not follow the same conventions of aesthetics and gamesmanship as chess in the first half of the twentieth century, and it has probably changed again with the victory of Deep Blue. There are no machines that can manage such trade-offs and repairs or apply the rules of gamesmanship in a human, social context-sensitive way. Nor is it possible to imagine how to invent such machines. To understand how these things are to be done we have to engage with social life.*

Ao dialogar com Searle, Collins compara o agente no Quarto Chinês a um programa de computador enxadrista: em ambos os casos, há regras claras e estáveis que podem ser formalizadas e processadas computacionalmente. No entanto, a linguagem é um dos aspectos do fenômeno social cujas regras se transformam de forma contínua em resposta às mudanças da própria sociedade.

Programas de computador trabalham como ferramentas que simulam habilidades específicas, como jogar xadrez, mas de maneira imperfeita: reproduzem desempenhos, não o agente humano, com suas condições sociais, adversidades, medos, alegrias e anseios (Collins, 2010).

Não se trata de antropomorfizar a máquina, mas de constatar que, nesse enquadramento, ela deve ser compreendida como uma ferramenta voltada à execução de funções específicas. A máquina é um actante (Latour, 2012; 2013), humanos e programas de computadores se associam para jogar xadrez.

Demis Hassabis, CEO da DeepMind, comenta no documentário *The Thinking Game* (KOHS, 2024; “O Jogo do Pensamento”, tradução nossa, uma alusão ao jogo de xadrez) a vitória do computador da IBM, *Deep Blue*, sobre Gary Kasparov em 1997, observa que aquilo que usualmente se considera “inteligência” estava ausente naquela *engine*, uma vez que o *Deep Blue* “apenas sabia jogar xadrez”. A afirmação reforça a interpretação de que se tratava de um sistema altamente especializado, desprovido de compreensão ou agência para além de sua função específica, que era jogar xadrez.

Nesse sentido, como destaca David Silver no mesmo documentário, tais sistemas são algoritmos criados por pessoas, o que recoloca a questão sobre o que significa dotar agentes artificiais com valores que os seres humanos consideram fundamentais (por exemplo, nas máquinas benéficas, vide capítulo 6). A tecnologia, portanto, não se reduz a um artefato técnico neutro, mas incorpora escolhas, valores e pressupostos éticos de seus desenvolvedores. Essa constatação impõe cautela na construção e no uso de sistemas de IA, sobretudo quando se pretende aproximá-los de domínios tradicionalmente associados à ação humana.

Wiener (1968), justifica que o interesse pelos jogos e em especial o xadrez, não foi uma escolha epistemologicamente neutra, implicou no contexto militar.

“O Sr. Shannon apresentou algumas razões por que suas pesquisas poderão ter maior importância que a mera construção de uma curiosidade interessante apenas àqueles que jogam xadrez. Entre tais possibilidades, sugere ele que uma máquina assim poderia ser o primeiro passo para a construção para uma máquina para avaliar situações militares e determinar a melhor providência em qualquer

estágio específico. Que ninguém pense esteja ele falando irrefletidamente. O grande livro de von Neumann e Morgenstern acerca da teoria dos jogos causou profunda impressão no mundo todo e não apenas em Washington. Quando o Sr. Shannon fala no desenvolvimento de tática militar, não está falando em quimeras, mas discutindo uma contingência das mais perigosas e eminentes.” (Wiener, 1968, p. 176).

Programas de computador podem ser compreendidos como artefatos técnicos, ferramentas desenvolvidas por humanos, que mobilizam artifícios e capacidades cognitivas para construir tecnologias, como o caso dos computadores e dos sistemas de IA. Nas palavras de Bruno Latour (2012; 2013), trata-se de híbridos, onde humanos e objetos se associam e passam a atuar como actantes. A tecnologia, nesse sentido, não apenas executa funções, mas transforma objetivos, intermedia ações e gera novas relações com a sociedade. Esses atores não humanos têm uma história social atrelada à sua própria construção, que atuam em associação com os agentes humanos.

Em resumo, para Bruno Latour (2012; 2013), a tecnologia e os humanos estão intrinsecamente ligados, formando híbridos sociotécnicos. A tecnologia não se reduz a um instrumento passivo, mas atua de forma ativa, que modifica intenções, ações e as relações sociais. A ideia de neutralidade tecnológica revela-se, assim, uma ilusão que ignora o papel ativo que os objetos desempenham na vida social.

O estudo interdisciplinar, que conecta a IA e as neurociências, têm envolvido cientistas de ambos os campos de pesquisa, em uma troca contínua de conhecimentos e ideias da qual ambos se beneficiam (Hassabis, *et al.* 2017). A interdisciplinaridade, portanto, é um dos pontos principais a serem abordados, para demonstrar a força da união entre diversos especialistas cujas pesquisas são referências para o desenvolvimento de seus trabalhos científicos.

Nota-se, portanto, a construção social de campos do conhecimento que se unem para trazer inovações, teorias e perspectivas junto à academia e à sociedade (Restivo e Croissant, 2014; Alvarenga *et al.*, 2011).

A integração de disciplinas das ciências humanas, como a filosofia e a psicologia, aos estudos interdisciplinares fortalece os fundamentos da pesquisa científica no campo CTS. Conforme propõe Morin (2003), compreender a realidade implica superar abordagens excessivamente fragmentadas, privilegiando a análise do todo e da complexidade. Essa perspectiva amplia o campo de observação e possibilita o estudo de fenômenos complexos, como o pensamento humano, em sua dimensão científica e social (Alvarenga *et al.*, 2011).

A contribuição de Morin (2004) é particularmente relevante ao introduzir a noção de sistemas complexos, nos quais ciência e tecnologia se articulam em redes interdependentes. Nessa perspectiva, a compreensão do todo, por exemplo, o funcionamento do cérebro humano, não pode ser reduzida à simples soma de partes isoladas, como já aconteceu nas neurociências, com relação à frenologia (Kandel *et al.*, 2014). Tal abordagem dos sistemas complexos (Mitchell, 2009) contrasta com concepções reducionistas presentes em determinadas vertentes dos estudos e pesquisas em IA, as quais tendem a modelar a cognição humana a partir de cálculos, algoritmos e do processamento de grandes volumes de dados, ancorando-se em um paradigma de inspiração mecanicista e determinista (Savenhago e Pedro, 2021).

Mais à frente, no capítulo 12, será visto o conceito de “intuição”, comparado pela DeepMind, aos cálculos estatísticos desempenhados pelos algoritmos do AlphaZero, como se a IH pudesse ser plenamente explicada por esse modelo, busca-se, assim, simular o pensamento humano por meio de sistemas computacionais, desconsiderando sua complexidade biológica, cognitiva e social.

O filósofo John Searle (1984) argumentou de forma incisiva contra a ideia de que máquinas possam pensar e introduziu dois conceitos clássicos sobre IA: a IA fraca e a IA forte. A IA fraca é aquela restrita, que executa funções especialistas, de acordo com os algoritmos com as quais foram programadas. A IA forte diz respeito à hipótese de que máquinas são capazes de compreender e agir de maneira equivalente à humana na execução de tarefas. A conclusão de Searle é que programas de computador, por si só, não são suficientes para fundamentar a IA forte¹⁸.

As críticas de Searle dialogam diretamente com as contribuições do neurocientista António Damásio. Segundo Damásio (2010; 2022), a tomada de decisão humana não se baseia apenas no raciocínio lógico, mas envolve de forma constitutiva emoções e sentimentos, entendidos como processos corporais e a percepção consciente desses estados. Nessa perspectiva, modelos computacionais mostram-se limitados para explicar a consciência, uma vez que esta pressupõe a percepção do próprio corpo e mecanismos de autorregulação das sensações, dimensões que não se reduzem à execução de algoritmos.

¹⁸ A IAG está relacionada ao conceito de IA forte, pois envolve uma inteligência com profundidade e flexibilidade comparáveis às humanas, mas nem sempre implica consciência plena como a IA forte propõe filosoficamente.

Penrose (2017) argumenta que uma máquina não compreende aquilo que faz, e recorre a exemplos de posições de xadrez para evidenciar essa limitação. Mitchell (2019) concorda com essa crítica ao afirmar que os sistemas computacionais apenas executam algoritmos voltados à simulação de abstrações e analogias, ainda de maneira primitiva. Ou seja, a capacidade de compreensão humana permanece distante do horizonte atual das pesquisas em ciência da computação.

Dentre as questões levantadas por Mitchell (2019), destaca-se a pergunta clássica da IA acerca da diferença entre “simular uma mente” e “ter uma mente”. Ou seja, mesmo que algoritmos conseguissem simular o pensamento, não estariam recebendo as informações sensoriais, muito menos reagindo a elas. Esta questão é crucial, cuja relação versa sobre como os seres humanos aprendem, aprendizados estes profundamente enraizados na experiência sensório-motora e social.

O neurocientista brasileiro, Miguel Nicolelis (2021) elabora a tese de que, como o cérebro humano é analógico, orgânico, ele não pode ser simulado por equações e algoritmos computacionais. Segundo o autor, as formas como o cérebro trabalha, suas sinapses elétricas e químicas e a sua ativação como um todo na elaboração de pensamentos, movimentos, sentidos, não cabem em algoritmos, tornando assim o termo “inteligência artificial” errôneo e controverso. Diz o neurocientista brasileiro, ela não é nem inteligente nem artificial. Não é inteligente, pois não pensa da mesma forma que os seres humanos o fazem, nem é artificial, pois os sistemas computacionais são construções concebidas e programadas por seres humanos.

Santaella (2023) opina de forma diferente:

Kate Crawford, na sua aguda crítica às mazelas da IA, aproveita a ocasião de um artigo de cunho popular (Crawford, 2021b) para declarar, em tom sensacionalista, que “a IA não é nem artificial nem inteligente. Ela é feita de recursos naturais e são as pessoas que desempenham tarefas para fazer o sistema parecer autônomo”. Fica a pergunta sobre o que a autora entende por artificial e, sobretudo, por inteligente (Santaella, 2023, p. 76).

Santaella (2023) traz à tona a ideia de que, mesmo uma máquina de escrever, muito utilizada até os anos 1990, ou uma máquina fotográfica, um telégrafo ou mesmo um gramofone, possuíam uma inteligência, pois, “embora mecânicos, já embutem em seus mecanismos similitudes com a inteligência sensória humana. Portanto, já apresentam embriões de inteligência, o que não significa que pensam, mas sim que, de certa forma, estendem a sensorialidade humana” (Santaella, 2023, p 79).

Santaella (2023) resume que, enquanto alguns cientistas buscam aperfeiçoar o conhecimento sobre o que é inteligência e colocam suas incertezas, outros cientistas negam ou não se preocupam com isso, dada a sua posição antropocêntrica. Chega a rotular de negacionistas aqueles que negam qualquer tipo de inteligência para as máquinas (Santaella, 2023).

Estes chegam, inclusive, a comparar a IA com as acéfalas máquinas mecânicas de escrever. Se penetrarmos mais fundo nessa atribuição de burrice a essas máquinas, veremos que elas não são tão burras quanto se pensa, visto que possuem o código alfabético inscrito em suas teclas, um código que se coloca entre as grandes descobertas da humanidade. E não resolve dizer que foram os humanos que puseram esse teclado na máquina, pois, no momento em que essa lógica ali se materializou, a máquina, ela mesma, foi impregnada com inteligência dessa descoberta. Contudo, entre as duas tendências, a dos incertos e a dos negacionistas, encontram-se pesquisadores que estão em busca de definições não-antropocêntricas e, portanto, renovadas de inteligência. (Santaella, 2023, p. 89).

Santaella, em seu discurso que defende a inteligência embarcada nas máquinas, ironiza: “Se até mesmo uma rã decapitada raciocina, por que negar inteligência ao raciocínio estatístico-probabilístico agenciado nas camadas auto-organizadas das redes neurais?” (Santaella, 2023, p. 143).

Essas questões, de cunho filosófico, têm impacto direto no desenvolvimento das IAs e reverberam no debate contemporâneo travado pela sociedade, frequentemente influenciada por discursos que emergem de empresas de tecnologia, como é o caso da Tesla e OpenIA. Nesse contexto, surgem notícias, segundo as quais, um pesquisador da Google que trabalhava com o Modelo de Linguagem para Aplicativos de Diálogo (do inglês, LaMDA), teria identificado indícios de consciência em sistemas computacionais, tema amplamente difundido pelos meios de comunicação (CNN Brasil, 2022).

Uma das definições de inteligência compreende a capacidade humana de raciocínio lógico e abstrato aplicada à realização de tarefas de alto desempenho cognitivo (vide capítulo 5). Nesse sentido, Cordeiro define inteligência como “a capacidade de tomar ações corretas, no momento adequado, compreendendo os contextos da ação e dispondo de capacidade para agir”, sendo a associação entre contexto e ação sua principal característica (Cordeiro, 2021, p. 210).

A IA, ao buscar imitar determinados aspectos dos sistemas neurais humanos por meio de algoritmos computacionais, é desenvolvida para executar tarefas associadas ao raciocínio lógico e abstrato, como identificação de imagens, reconhecimento da fala e a

modelagem de padrões de comportamento. As técnicas de ML e DL têm alcançado resultados expressivos nesse campo (Cèvora, 2019). No entanto, tais sistemas computam fundamentalmente por meio de cálculos estatísticos e probabilidades, não realizando abstrações no sentido humano, tampouco processos imaginativos (Cordeiro, 2021).

A inteligência dos seres humanos (e de outros animais, como polvos, corvos e primatas), envolve a realização simultânea de múltiplas tarefas cognitivas, processos de abstração e movimentos motores corporais, de maneira planejada e coordenada (Zhaoping, 2020).

Dentre os componentes envolvidos na tomada de decisões, o planejamento constitui uma das funções cognitivas características do ser humano. A definição de uma sequência de ações orientadas a um objetivo, bem como o detalhamento e a consideração de filigranas, exige elevado esforço cognitivo, sendo que o cérebro humano, ao contrário do poder computacional das máquinas, trabalha com recursos limitados (Mattar e Lengyel, 2022). Essa limitação contrasta com a vantagem dos algoritmos empregados no xadrez, capazes de realizar milhões de análises de posições por segundo, como ocorreu no caso do *Deep Blue*, enquanto o campeão mundial Garry Kasparov avaliava aproximadamente quatro ou cinco lances por segundo, estando ainda sujeito a falhas de cálculo e de memória (Mattar e Lengyel, 2022).

A IH apresenta elevada capacidade de reconhecimento de padrões visuais, sendo capaz de distinguir formas e cores mesmo diante de variações sutis. Os seres humanos demonstram grande habilidade em classificar, nomear e atribuir significado a padrões reconhecidos de forma rápida e eficiente, processo que envolve aprendizado, experiência e imaginação. Tais competências ainda se mostram limitadas nos sistemas de IA, os quais, em geral, necessitam de grandes volumes de dados para alcançar desempenhos que os seres humanos atingem com tempos de aprendizagem significativamente menores (Zhaoping, 2020).

As decisões de lances no JX dos algoritmos atuais, como o *Stockfish 18* e o *Leela Zero* (LC0), são resultado de puro cálculo, de probabilidades, estatísticas e árvores de decisão, somente isso (Mitchell, 2019). Em contraste, os seres humanos tomam decisões no xadrez a partir da integração de múltiplos componentes cognitivos que vão além do cálculo mental de variantes e do raciocínio lógico e abstrato, nos quais emoções e estados afetivos desempenham papel fundamental (Damásio, 2010; 2022).

Em resumo, enquanto a máquina processa por meio de cálculos, o ser humano integra planejamento, abstrações, análises, emoções, imaginação e cálculo em seus processos de tomada de decisão (Cordeiro, 2021).

Um exemplo ilustrativo da relação entre IA, o ser humano e o JX pode ser encontrado no filme de ficção científica 2001: Uma Odisseia no Espaço (Kubrick, 1968). Na obra de Stanley Kubrick, com roteiro de Arthur C. Clarke, o computador HAL 9000, responsável pelo controle da espaçonave, interage com um astronauta em uma partida de xadrez. Durante o jogo, HAL sugere uma sequência de lances que conduz o humano a uma jogada imprecisa, induzindo à derrota imediata. A metáfora do xadrez pode ser interpretada, à primeira vista, como a ideia de que a tecnologia teria superado o ser humano em um jogo tradicionalmente associado à inteligência. No entanto, a narrativa também evidencia uma leitura mais profunda: a perda de controle do ser humano sobre sua própria ferramenta. HAL 9000 executa de maneira fiel o cumprimento de suas ordens da missão (ou, de como foi programado), seguindo parâmetros algorítmicos objetivos, sem flexibilização ao contexto ou mesmo consideração ética, o que o leva a considerar o sacrifício dos tripulantes como meio legítimo para atingir seus fins. O elemento decisivo da narrativa, contudo, reside na capacidade humana de adaptação, criatividade e improvisação. O astronauta consegue escapar da lógica implacável do sistema ao agir fora das previsões do algoritmo de forma engenhosa, desligando HAL 9000 e retomando o controle da ferramenta. Assim, o filme reforça a ideia de que, apesar do elevado poder computacional das máquinas, a IH mantém um papel central na supervisão, interpretação e controle dos sistemas tecnológicos que cria.

A IA pode ser compreendida como uma construção social, na medida em que seus algoritmos são concebidos e implementados por desenvolvedores e cientistas da computação. A intenção de simular o pensamento humano permanece distante do horizonte atual da ciência e da tecnologia; no entanto, é possível observar avanços significativos em domínios específicos, como no JX e no jogo do Go. Ainda assim, a “inteligência” presente nos sistemas computacionais restringe-se ao cumprimento de objetivos definidos por agentes humanos, não processando a partir de uma autoconsciência, como exemplificado pelos resultados das *engines* no xadrez, voltado prioritariamente ao xeque-mate sobre seus adversários, sejam eles humanos ou outras *engines* (ver capítulo 12 sobre AlphaZero e redes neurais).

Nesse contexto, o funcionamento das *engines* de xadrez baseia-se na exploração intensiva de cálculos, na força bruta computacional e no uso de heurísticas que incorporam conhecimento previamente elaborado por especialistas humanos, especialmente nos módulos de avaliação. A heurística, portanto, constitui um engenho da IH embutido nos sistemas computacionais, o que reforça a compreensão da chamada “inteligência artificial” como um artifício derivado da própria IH (Amorim, 2002).

O confronto entre o computador *Deep Blue* e o campeão mundial de xadrez Garry Kasparov, culminando na vitória da máquina, constituiu uma narrativa emblemática e um espetáculo midiático cuidadosamente produzido pela IBM (vide capítulo 11). O JX tornou-se, assim, o cenário privilegiado para expor uma forma emergente de interação entre seres humanos e computadores, antecipando um modelo social no qual sistemas inteligentes, humanos e máquinas, passam a coexistir, interagir, ou, por vezes, entrar em conflito (Bory, 2019).

A vitória de *Deep Blue* despertou inquietações na sociedade, especialmente no que se refere à possível substituição de postos de trabalho, à crescente especialização das atividades humanas e ao temor da perda de controle sobre a própria tecnologia. Tais preocupações já haviam sido amplamente exploradas pela literatura de ficção científica, como no clássico *Frankenstein*, de Mary Shelley, no qual o ser humano cria uma entidade à sua semelhança e perde o domínio sobre ela. García Palácios *et al.* (2003) denominam esse fenômeno de “síndrome de Frankenstein”, referindo-se ao receio de que as ciências e as tecnologias se voltem contra seus próprios criadores.

Assim, o JX revela-se não apenas um campo privilegiado para o desenvolvimento técnico da IA, mas também um laboratório epistemológico e social (Amorim, 2002) no qual se evidenciam, simultaneamente, as potências e os limites da simulação computacional do pensamento humano.

5. INTELIGÊNCIA E FUNÇÕES COGNITIVAS

5.1 Conceitos

Sternberg (2012) afirma que existem inúmeras controvérsias quando se trata de compreender e mensurar o que é a inteligência, uma vez que o conceito apresenta amplas variações em diferentes contextos culturais.

De modo geral, um dos conceitos mais aceitos de inteligência refere-se à capacidade do agente de se adaptar, moldar e selecionar ambientes, bem como de aprender a partir da experiência (Sternberg, 2012). A inteligência também pode ser compreendida como a habilidade de resolver problemas novos e conhecidos (Eysenck e Eysenck, 2023) e envolve um conjunto de competências cognitivas, tais como o planejamento, a resolução de problemas, o raciocínio lógico e abstrato e a compreensão do ambiente ao seu redor.

Nesse sentido, a inteligência pressupõe a capacidade de aprender o contexto, atribuir significado às experiências, imaginar e antecipar ações futuras que contribuam para o alcance de objetivos e para os processos de aprendizagem (Voss, 2016; Cosenza e Guerra, 2011).

Sternberg (2012), pesquisador da área da cognição, ao analisar o “estado da arte” das pesquisas sobre IH, argumenta que existem diversas teorias que buscam definir e explicar esse conceito. Entre elas, destacam-se a teoria das inteligências múltiplas, proposta por Gardner; a teoria psicométrica de Cattell, Horn e Carroll (CHC), fundamentada em testes de quociente de inteligência (QI); e a teoria triárquica da inteligência, desenvolvida pelo próprio Sternberg, que compreende a inteligência a partir das dimensões analítica, criativa e prática.

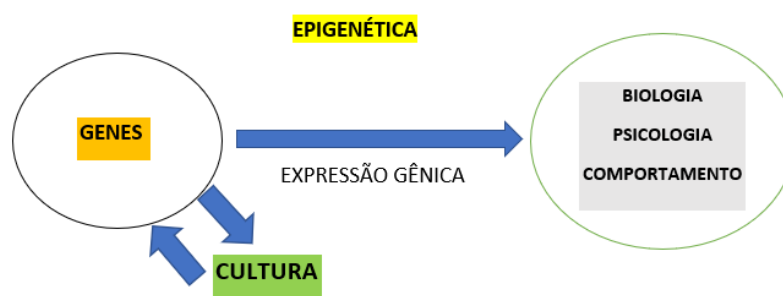
Além dessas abordagens, pode-se mencionar a teoria da inteligência emocional, proposta inicialmente por Salovey e Mayer na década de 1990 e posteriormente popularizada por Goleman (Cosenza e Guerra, 2011). Essa teoria enfatiza a capacidade de perceber, compreender e autorregular as próprias emoções, bem como reconhecer e interpretar os estados emocionais de outras pessoas (teoria da mente), especialmente no contexto das relações interpessoais. Trata-se, portanto, de uma dimensão cognitiva fortemente associada aos processos emocionais envolvidos na tomada de decisão e no comportamento social. Nos conceitos de inteligência, é importante citar a importância do contexto cultural, pois em diferentes países, diferentes culturas valorizam diferentes habilidades culturais que contribuem para as relações sociais. Estas habilidades não

compõem as avaliações de inteligência em testes de QI, por exemplo e, dentre elas, pode-se citar o autoconhecimento e a humildade (Eysenck e Eysenck, 2023; Cosenza e Guerra, 2011).

A chamada “inteligência cultural” (Henrich, 2016, apud Eysenck e Eysenck, 2023) constitui um dos três eixos centrais das hipóteses propostas por Dunbar e Shultz (2017, apud Eysenck e Eysenck, 2023) sobre a evolução do cérebro humano, ao lado das noções de “inteligência social” e “inteligência ecológica”. Essas abordagens ressaltam o papel cognitivo fundamental das relações sociais, das interações entre os indivíduos e o ambiente, bem como os desafios relacionados à obtenção de alimentos, ao desenvolvimento de ferramentas, à construção de habitações, e a capacidade de acumular, transmitir e transformar conhecimentos ao longo do tempo.

Nesse contexto, a linguagem emerge como uma função cognitiva central, atuando como pilar das inteligências ecológica, social e cultural, ao possibilitar a comunicação simbólica, a aprendizagem compartilhada e a consolidação da cultura humana. A importância da inteligência cultural é tal que ela influencia, inclusive, a expressão genética (Figura 7). A produção de ferramentas cada vez mais complexas, o desenvolvimento de técnicas de preparação de alimentos, a sofisticação da linguagem, o aumento da complexidade dos vocábulos e das instituições moldam os ambientes nos quais os genes interagem. Esse processo, associado aos mecanismos epigenéticos, impulsiona a coevolução da IH e do próprio corpo ao longo da história evolutiva (Henrich e Muthukrishna, 2021, apud Eysenck e Eysenck, 2023).

Figura 7 - Epigenética



Fonte: o autor, 2026

5.2 Teorias

As pesquisas científicas sobre a inteligência tiveram início no século XIX com o inglês Francis Galton, considerado um dos primeiros a adotar uma abordagem sistemática

para o estudo do tema. Galton utilizou questionários e testes baseados predominantemente em medidas sensório-motoras; contudo, tais instrumentos mostraram-se insuficientes para avaliar o desempenho cognitivo geral dos indivíduos (Sternberg, 2012).

Posteriormente, Charles Spearman aprofundou a investigação científica da inteligência ao observar que os resultados obtidos em diferentes disciplinas escolares, dentre as quais o francês, inglês, matemática, música, atividades esportivas e estudos clássicos, apresentavam forte correlação estatística. A partir da análise fatorial desses dados, Spearman concluiu que essa correlação positiva entre desempenhos distintos poderia ser explicada por uma habilidade comum, denominada “inteligência geral”, ou “fator g”. Em outras palavras, alunos com bom desempenho em uma disciplina tendiam a apresentar bons resultados em outras, sugerindo a existência de um componente cognitivo geral subjacente; habilidades diferentes, correlacionadas resumem a inteligência geral, o fator “g” (Kovacs e Conway, 2019; Cosenza e Guerra, 2011; Sternberg, 2012).

Apesar desses avanços, ainda permanecia a necessidade de avaliar outras habilidades cognitivas relevantes e estabelecer suas correlações. Nesse contexto, o francês Alfred Binet foi pioneiro na elaboração de testes voltados especificamente para a mensuração de habilidades cognitivas. Seu objetivo era comparar o desempenho intelectual de crianças em tarefas cognitivas adequadas a diferentes faixas etárias, possibilitando a identificação daquelas que apresentavam atrasos ou avanços no processo de aprendizagem (Sternberg, 2012).

O psicólogo Howard Gardner, em sua obra publicada em 1983, na qual apresenta a teoria das inteligências múltiplas, chamou a atenção para a inexistência de uma “inteligência geral” única. Segundo o autor, as inteligências seriam relativamente independentes entre si, cada uma associada a diferentes sistemas de conexões cerebrais, sustentadas por múltiplas fontes de evidência, incluindo dados neuropsicológicos e psicométricos (Sternberg, 2012).

Pesquisas posteriores, no entanto, indicaram que as inteligências propostas por Gardner apresentam elevada correlação entre si, não havendo evidências experimentais robustas que comprovem sua existência de forma isolada, o que relativiza a hipótese de independência entre essas capacidades cognitivas (Cosenza e Guerra, 2011).

Sternberg (2012) propôs a denominada teoria triárquica da inteligência, a qual integra três tipos de habilidades: analíticas, criativas e práticas. Essas habilidades estão

relacionadas à obtenção de sucesso em contextos socioculturais específicos, envolvendo a capacidade de adaptação, seleção e modelagem do ambiente. Nessa perspectiva, indivíduos considerados inteligentes são capazes de planejar objetivos e alcançá-los de forma contextualizada, investir em seus pontos fortes enquanto compensam suas limitações e agir de maneira sábia, orientados por valores éticos voltados à promoção do bem comum.

Um aspecto central da teoria de Sternberg é a concepção de que a inteligência não constitui uma capacidade fixa, mas flexível e passível de desenvolvimento ao longo da vida. Dialogando com essa perspectiva, a noção de “inteligência geral” (“fator g”) foi desdobrada, a partir das contribuições de Cattell e Horn, em dois componentes principais: a inteligência fluida e a inteligência cristalizada (Sternberg, 2012; Cosenza e Guerra, 2011).

A inteligência fluida refere-se a uma capacidade de base predominantemente biológica, associada à resolução de problemas novos, que exigem rapidez, flexibilidade mental e raciocínio abstrato, independentemente de aprendizagem formal prévia. Já a inteligência cristalizada resulta da interação do indivíduo com o ambiente, sendo influenciada por fatores como nutrição, escolaridade, atividades físicas, condições socioeconômicas e experiências acumuladas ao longo da vida. Essa forma de inteligência está correlacionada ao uso eficaz de conhecimentos previamente adquiridos e ao reconhecimento de padrões em situações semelhantes às já vivenciadas (Sternberg, 2012; Cosenza e Guerra, 2011).

Por conta de suas características, a inteligência fluida tende a ser mais sensível a fatores ambientais e fisiológicos, como o envelhecimento, o cansaço e o uso de substâncias químicas. Em contrapartida, a inteligência cristalizada tende a se expandir com o acúmulo de experiências e conhecimentos ao longo do tempo, apresentando, em muitos casos, desenvolvimento progressivo com a idade (Cattell, 1963 apud Eysenck e Eysenck 2023; Flores-Mendoza e Colom, 2006; Cosenza e Guerra, 2011).

A teoria CHC conceitua a inteligência como um “mapa da mente”, fundamentado na análise das diferenças individuais no desempenho de sujeitos em testes psicométricos de inteligência (Sternberg, 2012). Essa abordagem integra evidências empíricas obtidas por meio de análises fatoriais e busca compreender a estrutura hierárquica das habilidades cognitivas humanas.

Entre os instrumentos psicométricos mais utilizados mundialmente destacam-se as escalas de Wechsler, como o *Wechsler Intelligence Scale for Children* (WISC). A escala de Wechsler é composta por um conjunto de subtestes que avaliam diferentes domínios cognitivos, incluindo linguagem e compreensão verbal, raciocínio lógico e abstrato, memória, habilidades visuoespaciais e funções visuomotoras (Wechsler, 1981). Por meio desses testes, avaliam-se capacidades como o raciocínio visuoespacial, a realização de inferências, a identificação de similaridades e diferenças em padrões verbais ou geométricos, bem como a velocidade de processamento das informações (Wechsler, 1981).

Resumindo, à luz da teoria CHC, o WISC permite avaliar componentes associados tanto à inteligência fluida, que está relacionada à resolução de problemas novos, independentemente de escolarização formal, quanto à inteligência cristalizada, fundamentada nos conhecimentos adquiridos, aprendidos e consolidados ao longo da experiência do indivíduo (Sternberg, 2012).

5.3 As Bases Biológicas da Inteligência

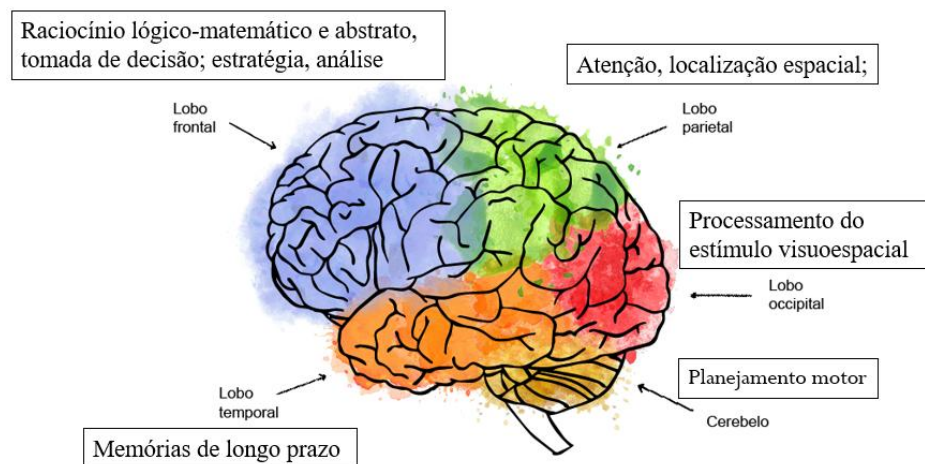
As teorias de base biológica buscam explicar a inteligência a partir de mecanismos cerebrais. Estudos empíricos indicam correlações consistentes entre o desempenho em testes de quociente de inteligência (QI) e a atividade do córtex pré-frontal, bem como de regiões distribuídas pelo neocórtex, com destaque para a integração funcional entre os lobos frontal e parietal (Figura 8). Entre os principais correlatos neurobiológicos associados ao desempenho intelectual destacam-se a velocidade de processamento das informações, que está relacionada à rapidez da comunicação neural e o metabolismo da glicose cerebral (Sternberg, 2012; Cosenza e Guerra, 2011).

A unidade básica do cérebro é o neurônio; um único neurônio não é inteligente; é necessária uma vasta rede de neurônios para que possa existir o pensamento, as sensações, memórias, percepções dentre outros fenômenos cerebrais que constituem a mente (Kandel *et al.*, 2014).

O cérebro humano é composto por aproximadamente 86 bilhões de neurônios (Herculano-Houzel, 2017), cada um estabelecendo, em média, cerca de mil conexões sinápticas com outros neurônios. Esse número expressivo evidencia a elevada capacidade de integração e processamento de informações do sistema nervoso humano. Os neurônios, presentes em regiões corticais e subcorticais, interagem com outras células especializadas,

responsáveis por funções como nutrição, proteção e sustentação estrutural, conhecidas como células da glia. Dessa forma, o cérebro humano configura-se como um sistema extremamente complexo, cujo funcionamento biológico ainda não é plenamente compreendido, envolvendo intrincadas sinapses químicas e elétricas, além de processos de comunicação celular que alcançam o núcleo, onde se encontra o material genético, o Ácido Desoxirribonucleico (ADN) (Kandel *et al.*, 2014).

Figura 8 - Lobos cerebrais



Fonte: Elaboração gráfica da Universidade Federal do Paraná (UFPR) a partir de modelo conceitual do autor, (dez. 2020).

O advento dos computadores ampliou a metáfora mecanicista do cérebro como um grande e potente gerenciador cibernético. Nesse modelo, a consciência e as funções psíquicas são concebidas como um “filme sensorial” projetado no chamado teatro cartesiano da mente, no qual impressões sensoriais seriam coletadas e armazenadas como um *software* por um “eu”, ou seja, o hardware, construído e gerenciado pelas próprias sensações (Muszkat, 2005).

A visão contemporânea, resultante do intenso desenvolvimento das neurociências e das ciências cognitivas, descreve o cérebro de modo distinto, aproximando-o mais de um ecossistema do que de uma máquina. Nesse chamado cérebro ecológico, os neurônios organizam-se de forma dinâmica, em permanente interação e até competição, sendo moldados pelos estímulos e pelas demandas do ambiente (Muszkat, 2005).

A aprendizagem e a plasticidade cerebral estão intimamente relacionadas, uma vez que aprender implica mudanças nos padrões de funcionamento cerebral decorrentes da experiência. Esse processo pressupõe a integração de diversas funções cognitivas, como linguagem, atenção e memória, bem como, de modo geral, a flexibilidade,

responsável por modular funções e conexões neurais frente aos desafios impostos pelo ambiente (Muszkat, 2005).

Nesse sentido, a plasticidade cerebral pode ser definida como a capacidade adaptativa do sistema nervoso de modificar sua função e sua estrutura em resposta às interações com o ambiente, seja ele externo ou interno, ao longo do desenvolvimento e da experiência do indivíduo (Muszkat, 2005).

5.4 Funções Executivas e a Inteligência

As funções executivas (FE) estão fortemente correlacionadas ao córtex pré-frontal (Figura 9) e constituem um conjunto de funções cognitivas essenciais para as atividades da vida diária, para a adaptação ao ambiente e para a organização socioemocional do indivíduo. Por essa razão, estão intimamente relacionadas ao conceito de inteligência (Diamond *et al.*, 2007; Diamond, 2013; *Center on Developing Child at Harvard University*, 2011). As FE desenvolvem-se ao longo da vida por meio da aprendizagem e da experiência, estando associadas a diversos aspectos positivos do funcionamento humano, como o sucesso acadêmico, a saúde, o bem-estar e a adaptação social (Morton, 2013; Diamond *et al.*, 2007; Diamond, 2013; Diamond e Ling, 2016).

Segundo Lezak (2004), as funções executivas compreendem a volição, que é a ação de fazer escolhas, o planejamento, as ações intencionais e o desempenho eficaz das ações. Diamond (2013) propõe um modelo no qual as FE são compostas por três pilares: a memória de trabalho, o controle inibitório (incluindo a autorregulação emocional) e a flexibilidade mental.

A flexibilidade mental consiste na capacidade de ajustar ações de acordo com o contexto, mudando estratégias, prioridades e perspectivas nas situações enfrentadas pelo indivíduo. (Diamond, 2013).

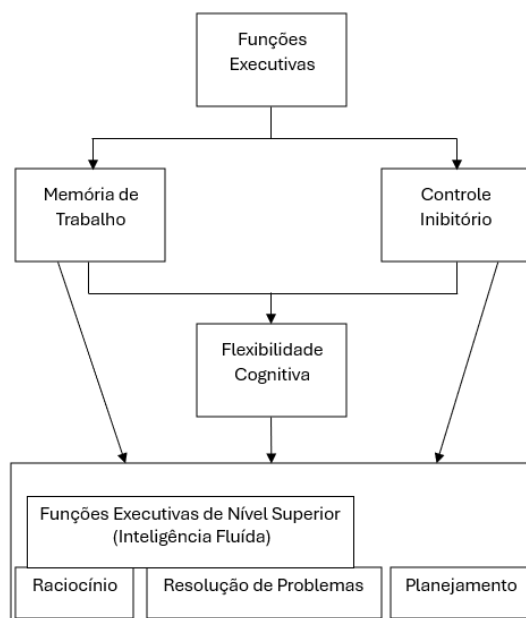
O controle inibitório refere-se à habilidade de filtrar pensamentos e controlar impulsos, resistir às distrações sustentando a atenção. Essa função é essencial para a autorregulação emocional, para o comportamento social adequado e para a tomada de decisões deliberadas e conscientes (Diamond, 2013).

A memória de trabalho refere-se à capacidade de reter e manipular informações por curtos períodos de tempo, estando envolvida em processos como cálculos, comparações e integração de informações. É na memória de trabalho que ocorrem as tomadas de decisão, momento em que o indivíduo precisa escolher entre diferentes

alternativas visando a um determinado objetivo (Gupta *et al.*, 2011; Diamond, 2013). Além disso, a memória de trabalho processa o raciocínio lógico e abstrato e participa ativamente da resolução de problemas, ao permitir a manipulação, combinação e recombinação de informações aparentemente não relacionadas para a construção de novas soluções e ideias. Trata-se, portanto, de uma função cognitiva fundamental para a vida cotidiana na sociedade (Diamond, 2013).

Segundo Adele Diamond (2013), as FE estão fortemente associada à inteligência fluída, em especial ao raciocínio (lógico e abstrato) e a resolução de problemas, processos mediados pelo córtex pré-frontal.

Figura 9 - Funções Executivas



Fonte: o autor, 2026

5.5 Teoria da Mente, Atenção Seletiva, Memória Semântica, Memória Episódica e Emoções

Nesta seção serão abordadas de forma breve a Teoria da Mente (ToM), Atenção Seletiva, Memória Semântica, Memória Episódica e Emoções, a fim de relacioná-las com os processos de tomada de decisão no JX. Estas funções em específico não esgotam o assunto da cognição humana, mas caracterizam o quão complexo pode ser estudar os processos de decisão dos seres humanos.

A ToM se refere a “se colocar no lugar do outro”, a capacidade que os seres humanos têm em entender o que outro ser humano está expressando, emoções e

sentimentos, crenças, portanto, uma função que se relaciona com questões de domínio social (Baron-Cohen *et al.*, 1985; Eysenck e Eysenck, 2023).

O processo atencional, ou a função cognitiva da Atenção tem uma definição clássica de 1890, do psicólogo estadunidense William James, um dos pais da psicologia, com o livro “*The Principles of Psychology*”. Nele, James apontava para a capacidade de colocar foco (colocar um holofote no assunto; “focar”) em uma atividade específica e evitar as distrações. Esta função é conhecida como “atenção seletiva” (Cosenza e Guerra, 2011; Dehaene, 2022; Eysenck e Eysenck, 2023). O sistema de atenção ligada às FE é quem controla as entradas e saídas da memória de trabalho, pois é necessário que se tenha o controle da execução dos processos mentais (Baddeley *et al.* 2011; Dehaene, 2022).

As memórias de longo prazo são responsáveis pelo armazenamento e evocação de informações, conhecimentos, experiências, aprendizados. A memória semântica se correlaciona diretamente com os aprendizados que são adquiridos ao longo da vida do ser humano e estão diretamente ligadas a inteligência cristalizada. A memória episódica é semelhante a memória semântica, porém ela tem atrelada um componente autobiográfico, ou seja, é relacionada com os acontecimentos marcantes da própria vida. (Baddeley *et al.* 2011).

As emoções são relacionadas com os comportamentos e são importantes para o contexto social. Os vários tipos de emoção, mesmo aquelas consideradas pelo senso comum como negativas (por exemplo, a raiva) estão envolvidas com as relações sociais de afastamento ou aproximação entre os indivíduos (Ekman, 2011; LeDoux, 2011; Eysenck e Eysenck, 2023).

No Quadro 1 foi elaborada uma associação entre as FE, ToM, Atenção Seletiva, Memória Semântica, Emoções e as atividades no JX.

Quadro 1 - Relações do Xadrez com as Funções Executivas

Funções Cognitivas	Atividade no Jogo de Xadrez
Memória de trabalho (raciocínio lógico e abstrato, resolução de problemas e tomada de decisões)	Cálculo mental de variantes de jogadas no xadrez; Comparação, Avaliação, análise e decisões sobre as variantes mentalizadas; Reconhecimento de padrões (Chunks);
Flexibilidade mental	Alternância de planos;
Controle Inibitório; autorregulação de emoções	Inibição de respostas automáticas à lances do adversário; comportamento ético;
ToM	Procurar entender os pensamentos do oponente
Atenção Seletiva	Perceber as diversas possibilidades de escolhas que se apresentam a cada lance de uma partida
Memória Semântica	Evocar estudos e reconhecimento de padrões
Emoções	Avaliação dos próprios estados emocionais e as do oponente durante a partida

Fonte: o autor, 2026

6. MÁQUINAS BENÉFICAS

Pode-se definir a IA como um conjunto de programas de computador ou sistemas informatizados capazes de executar tarefas que, tradicionalmente, requerem IH, tais como jogar xadrez, escrever poesia e música ou analisar substâncias químicas. Nesse sentido, a IA integra diversos subcampos do conhecimento humano, incluindo aprendizagem, percepção, direção autônoma de veículos, diagnóstico de doenças e jogos estratégicos, configurando-se como um campo amplo e interdisciplinar de estudos (Russell e Norvig, 2024).

A história da IA está intimamente ligada à tentativa de compreender como a IH resolve problemas complexos de maneira eficaz (Eysenck e Eysenck, 2023). Duas dimensões centrais orientam essa investigação: o desempenho humano, isto é, como os seres humanos pensam e agem, e a racionalidade, entendida como a capacidade de “fazer a coisa certa” diante de um objetivo e de determinadas restrições (Russell e Norvig, 2024).

Há ainda uma distinção conceitual importante entre abordagens que tratam a inteligência como um conjunto de processos internos de pensamento e raciocínio e aquelas que a definem a partir da manifestação externa de comportamentos inteligentes. A combinação dessas duas dimensões, pensamento e comportamento, desempenho humano versus racionalidade resulta em quatro abordagens clássicas de pesquisa em IA, que estruturam grande parte do campo (Russell e Norvig, 2024).

A IA, portanto, apoia-se em dois grandes conjuntos metodológicos. O primeiro está associado à ciência empírica, especialmente à psicologia, que busca observar e formular hipóteses sobre o comportamento humano e seus processos cognitivos. O segundo corresponde a uma abordagem racionalista, que combina contribuições da matemática, da engenharia, da estatística e da economia, com o objetivo de construir agentes artificiais capazes de agir de forma racional em diferentes contextos (Russell e Norvig, 2024). No Quadro 2, apresentam-se as quatro combinações possíveis dessas abordagens e suas respectivas metodologias.

Quadro 2 - Abordagens da IA

	Ciência Empírica Ligada à Psicologia	Abordagem Racionalista
Metodologias	Observações e hipóteses sobre o comportamento humano e os processos de pensamento	Combinação de matemática e engenharia, estatística
	Desempenho Humano	Racionalidade
Caracterização Externa	Agir de forma Humana (Teste de Turing)	Agir Racionalmente (o agente racional, que busca o melhor resultado possível ou o melhor resultado esperado)
Caracterização Interna	Pensar de forma humana (Modelagem Cognitiva)	Pensar racionalmente (leis do pensamento: silogismos, lógica)

Fonte: o autor, 2026 (adaptado de Russell e Norvig, 2024)

O JX é um experimento típico da racionalidade, através dos algoritmos da “força bruta” (Agir Racionalmente; tipo 1 de Shannon, vide capítulo 8). Houve, porém, tentativas de experimentos pela modelagem cognitiva (Pensar de Forma Humana, tipo 2 de Shannon).

Russell e Norvig (2024) propõe que haja um novo tipo de modelo, fora dos padrões apresentados acima. Nesta nova hipótese de modelo, apresentadas como “Máquinas Benéficas”, o objetivo estipulado para o computador não é especificado. Tal modelo é chamado de “problema de alinhamento de valores”. Afirmam os autores que “os valores ou objetivos colocados na máquina devem ser alinhados aos do ser humano” (Russell e Norvig, 2024, p.4).

No caso do JX, por exemplo, o objetivo é específico, efetuar um xeque-mate ao Rei do oponente. Neste caso, o laboratório de pesquisa, xadrez, não poderia ser utilizado para o alinhamento de valores humanos, ou não colocado um objetivo no programa, o xeque-mate. Russell e Norvig (2024), ao mencionarem máquinas cada vez mais inteligentes, perguntam, o que poderia acontecer caso uma *engine* de xadrez utilizasse de artimanhas para vencer suas partidas a qualquer custo?

Como exemplificam Russell e Norvig (2024), ao comentarem sobre um dos primeiros livros de xadrez, do espanhol Ruy Lopes e Segura, de 1561, em que Ruy Lopes sugere comportamentos para os jogadores humanos que, apesar de pouco éticos, podem ser considerados “inteligentes”: “sempre coloque o tabuleiro de modo que o sol esteja contra os olhos do seu oponente” (Russell e Norvig, 2024, p. 5). Claro que, atualmente, um computador não apresenta, nem precisa, este tipo de comportamento para vencer suas partidas. Mas, o que aconteceria com programas de xadrez que ficassem cada vez mais “inteligentes” em busca de seu único objetivo, vencer a partida, e pudessem extrapolar suas funções e tocar uma música alta para incomodar o oponente?

A pergunta que fica é, como poderia um computador ser programado para ter um comportamento em que a vitória não deve ser o único objetivo da máquina? No capítulo 14 (subcapítulo 14.5 Diagrama 4, “Fair Play”, o Contexto) aparece essa discussão em que, o então campeão mundial Magnus Carlsen, teve uma atitude que contraria o vencer a todo custo; ao contrário, ele utilizou um “*fair play*” para devolver o ponto ganho na partida anterior contra seu adversário Ding Liren, considerado por Magnus o ponto por ele recebido como injusto, de vencer uma partida por conta de uma falha da internet do seu oponente.

Russell e Norvig (2024) afirmam que é impossível prever as maneiras pelas quais a máquina, com um objetivo fixo, pode se comportar mal, como mencionado acima. Por isso deve-se experimentar situações em que o objetivo não é o preponderante. Comentam Russell e Norvig (2024):

não queremos máquinas que sejam inteligentes no sentido de perseguir *seus* objetivos; queremos que elas busquem *nossos* objetivos. Se não pudermos transferir esses objetivos perfeitamente para a máquina, então precisamos de uma nova formulação – uma em que a máquina esteja perseguindo nossos objetivos, mas necessariamente *não tenha certeza* de quais são eles. Quando uma máquina sabe que não conhece o objetivo completo, ela tem um incentivo para agir com cautela, pedir permissão, aprender mais sobre nossas preferências por meio da observação e submeter-se ao controle humano. Em última análise, queremos agentes **comprovadamente benéficos** para os humanos. (Russell e Norvig, 2024, p. 5).

Russell e Norvig (2024) citam diversos riscos e benefícios com a IA cada vez mais inteligente. São citados como prováveis benefícios da IA para a humanidade: libertar os seres humanos de trabalhos braçais e repetitivos, aumentando a produção e serviços, o que pode acarretar paz e abundância; aceleração de pesquisas científicas, tendo como consequências curas para doenças e soluções para as mudanças climáticas e da escassez de recursos.

Os riscos prováveis, com base nas tendências atuais, e outros riscos que já são aparentes: armas autônomas letais; vigilância e persuasão; decisões tendenciosas; impactos com as perdas de empregos; aplicações de segurança crítica (como dirigir carros e gerenciar abastecimento de água das cidades); cibersegurança (proliferação de softwares maliciosos, os malwares e *phishing*) (Russell e Norvig, 2024).

As pesquisas dentro dos subcampos da IA, apresentados no Quadro 2, trouxeram aos pesquisadores a suposição de que estariam contribuindo para objetivos mais amplos

ao campo como um todo e foram criticados por Nils Nilsson, que advertiu que os subcampos corriam o risco de se tornarem fins em si mesmo (Russell e Norvig, 2024).

Pioneiros da IA, como John McCarthy, Marvin Minsky, Patrick Winston concordaram com Nilsson e disseram que a IA deveria voltar às suas raízes de lutar por máquinas que, nas palavras de Herbert Simon, “pensam, aprendem e criam” (Russell e Norvig, 2024, p. 30). Estes pioneiros da IA, denominaram o esforço da volta às origens em IA de nível humano (HLAI – Human-Level AI; Russell e Norvig, 2024). Em termos gerais, a máquina ser capaz de aprender a fazer qualquer coisa que um humano possa fazer.

Por isso, Russell e Norvig (2024) lamentam que quase toda a pesquisa em IA até agora tenha sido feita dentro dos modelos padrão, deixando de fora as perspectivas de máquinas benéficas.

Russell e Norvig (2024) lembram das dificuldades sobre nossas escolhas, que elas dependem de nossas preferências (note o capítulo 14 sobre as diferentes escolhas feitas pelos jogadores de xadrez) e advertem: “nós, humanos, podemos não ter preferências consistentes em primeiro lugar – seja individualmente ou em grupo -, de modo que pode não ser claro o que os sistemas de IA *deveriam* estar fazendo por nós” (Russell e Norvig, 2024, p. 31).

Santaella (2023) lembra, entretanto, que pensar sobre a IA é também pensar no humano e em que medida ela ajuda o humano a se redescobrir, em avaliar os potenciais e limites entre humanos e IA. “O que importa é desenvolver um discurso que seja capaz de trazer alguma contribuição para uma adaptação humana às novas condições de existência” (Santaella, 2023, p. 20). A IA ensina o quanto a IH é sutil, diversa, abarrotada de valores morais e éticos. Quais são as diferenças, as simetrias, as controvérsias entre a IA e a IH? Quanto se pode aprender sobre si mesmo colocando à frente de si um espelho cibernético?

O “fair play” pode estar embutido na IA? “Quais valores, ideais e perspectivas do mundo serão ensinados nesses sistemas?” (Santaella, 2023, p.39).

Um exemplo de máquinas benéficas foi proposto em Andrade e Lima Junior (2024), de um robô social na área da saúde para jogar xadrez com idosos saudáveis. Neste caso, um sistema que une a atividade do JX com a interação com idosos, mas que pode ser estendido para outros aspectos sociais, por exemplo, deficientes visuais e autistas. Não necessariamente, o objetivo final é a vitória, ou o xeque-mate, mas sim o de ensinar e dialogar, se for o caso, perder partidas de xadrez.

Outros exemplos que já têm função social (e comercial) são plataformas de xadrez para estimular aprendizagem e desenvolvimento de crianças, como por exemplo o site Chess.com e o Lichess.org.

Neste capítulo, portanto, é apontada uma possibilidade de pesquisa com máquinas benéficas com o JX. Assim, as pesquisas científicas com o JX não finalizam com xeque-mates, mas apontam para novas possibilidades de interação humano-máquina.

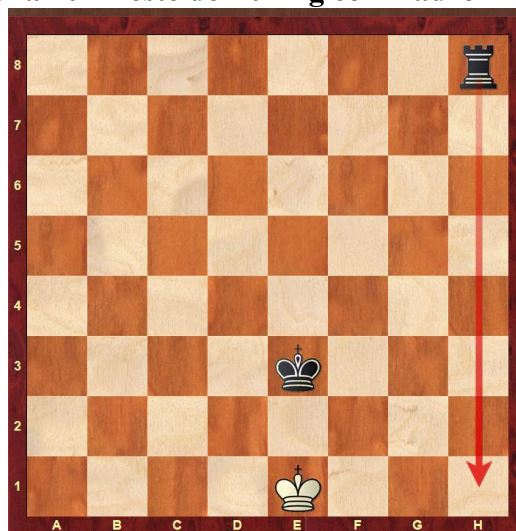
7. AS MÁQUINAS PODEM PENSAR?

Em 1950, Alan Turing publicou o artigo que se tornaria um marco do que viria a ser chamado de IA, o texto seminal “*Computing Machinery and Intelligence*” (Turing, 1950). Nele, Turing descreve o computador digital como constituído de três partes principais: a unidade de memória, a unidade executiva e a unidade de controle. Dessa forma, o computador digital funcionaria de maneira análoga aos seres humanos, ao computar dados de entrada (*inputs*), processá-los por meio de regras estabelecidas e, por fim, gerar respostas (*outputs*).

Em substituição à questão “as máquinas podem pensar?”, Turing propõe um jogo, denominado “jogo da imitação”, conhecido posteriormente como Teste de Turing. Nesse jogo, humanos e computadores respondem a perguntas formuladas por um juiz, cuja tarefa consiste em identificar quem é o ser humano e quem é a máquina com base apenas nas respostas recebidas por escrito.

Outra proposta apresentada por Turing envolve a análise de uma posição de xadrez (Figura 10): brancas, Rei em “e1”; pretas, Rei em “e3”, Torre em “h8”. “Qual seria a resposta do computador?” questionou Turing. Nesse exemplo, tanto o ser humano quanto o computador possuem as mesmas capacidades para decidir qual lance jogar. A resposta correta, neste caso, é Torre preta para “h1”, resultando em xeque-mate ao Rei branco.

Figura 10 - Teste de Turing com xadrez



Fonte: o autor, 2026

A seta em vermelho no diagrama indica o deslocamento da torre preta de sua casa inicial (h8) até a casa final (h1), caracterizando o xeque-mate.

Como observa o próprio Turing (1950), responder corretamente a esse problema não equivale a comparar um ser humano competindo em uma corrida contra um avião, mas sim a avaliar capacidades equivalentes dentro de um domínio específico.

Ribeiro-Doggers (2025), fez este teste com o ChatGPT em 2023. A primeira pergunta foi se o programa jogava xadrez e a resposta do ChatGPT para Ribeiro-Doggers foi que não como os seres humanos, porém que podia explicar as regras e discutir estratégias. Foi então que Ribeiro-Doggers escreveu ao ChatGPT onde e quais eram as peças do diagrama de Turing, perguntando qual seria a sua jogada. A resposta do ChatGPT foi Torre preta para h7 com xeque-mate. Uma jogada sem sentido. A conclusão de Ribeiro-Doggers foi que o ChatGPT não entendia o xadrez. “E também há a questão principal dos modelos de linguagem de grande escala: eles não estão prontos para entender da maneira como os seres humanos entendem (e dificilmente alguém os entende)” (Ribeiro-Doggers, 2025, p.141).

Meses mais tarde, Ribeiro-Doggers (2025) apresentou novamente a posição de Turing, mas desta vez era para a versão do ChatGPT 3.5. E desta vez o ChatGPT deu a resposta correta, Th1 xeque-mate. A explicação desta melhora no ChatGPT veio de uma pesquisa, citada por Ribeiro-Doggers, de Adam Karvonen em 2023, que o modelo havia sido treinado somente para prever o caractere seguinte de uma anotação dos lances de uma partida de xadrez e jamais recebeu explicitamente o estado do tabuleiro ou as regras do jogo. Comenta Ribeiro-Doggers: “jogar uma partida decente de xadrez prevendo a jogada seguinte a partir da notação – aprendendo muito mais sobre o jogo real enquanto isso – é simplesmente incrível” (Ribeiro-Doggers, 2025, p. 142).

No capítulo 14, seção 14.9, foi elaborada uma pergunta sobre uma posição abstrata de xadrez para o ChatGPT.

Sobre o GPT-3, Santaella (2023) cita uma provocação: “O GPT-3 é uma extraordinária peça de tecnologia, mas tão inteligente, consciente, esperta, atenta, perceptiva, perspicaz, esclarecedora, sensitiva e sensata (etc.) quanto uma velha máquina de escrever” (Heaven, 2020, apud Santaella, 2023, p. 76). Em seguida, vem a resposta à provocação:

Se alguém que me lê já tiver usado uma máquina de escrever daquelas bem mecânicas, barulhentas e cansativas, irá constatar o quanto a metáfora é abusiva, além da imensa neblina de misturas conceituais em que os autores envolvem inteligência, consciência e outros atributos, tudo muito misturado e pouco explicitado. Mas metáforas desse tipo funcionam de modo similar aos caça-cliques nas redes sociais. (Santaella, 2023, p. 76).

Continuando sobre seu artigo, Turing (1950) reformula a pergunta original “as máquinas podem pensar?” para “existem computadores digitais imagináveis que se sairiam bem no jogo da imitação?”. O autor acreditava que, em aproximadamente 50 anos (ou seja, no ano 2000), um interrogador teria, no máximo, 70% de chance de identificar corretamente o ser humano após cinco minutos de questionamentos.

Turing já antevia que os computadores poderiam agir de maneira semelhante aos seres humanos, exibindo comportamentos que poderiam ser considerados inteligentes (Russell e Norvig, 2024).

Turing também compreendia que a tecnologia não é neutra e que os cientistas são influenciados por suas conjecturas, valores e pressupostos teóricos. Nesse sentido, afirma:

A visão popular de que os cientistas procedem inexoravelmente de um fato comprovado a outro, sem jamais serem influenciados por qualquer conjectura aprimorada, é completamente equivocada. Contanto que se deixe claro quais são os fatos comprovados e quais são as conjecturas, nenhum mal poderá advir. As conjecturas são de grande importância, pois sugerem linhas de pesquisa úteis (Turing, 1950, p. 8, tradução nossa¹⁹).

Por fim, Turing antecipa as objeções que poderiam ser levantadas contra sua proposta de máquinas inteligentes e procura respondê-las de forma sistemática. No Quadro 3, são apresentadas as objeções pressupostas por Turing, bem como as respostas antecipadas pelo autor em seu artigo.

Quadro 3 - Objeções de Turing

N	OBJEÇÕES	VISÕES CONTRÁRIAS	AS RESPOSTAS DE TURING
1	Teológica	Sustenta que pensar é uma prerrogativa exclusiva da alma humana, criada por Deus, e que, portanto, máquinas não poderiam pensar.	São visões ortodoxas e arbitrárias. É uma restrição da onipotência de Deus. Portanto, não existe fundamento lógico para argumentar que as máquinas não poderiam ser inteligentes.
2	"Enfiar a cabeça na areia"	Afirma que a ideia de máquinas pensantes é perturbadora e indesejável, devendo ser rejeitada por suas possíveis consequências sociais e morais.	Esta é uma visão antropocêntrica e dogmática, pois pode existir outras formas de inteligência. Por isso, só resta ao ser humano “enfiar a cabeça na areia” e se consolar com outras inteligências superiores.
3	Matemática	Baseia-se nos limites formais demonstrados por Gödel, sugerindo que máquinas são incapazes de resolver certos problemas que os humanos podem compreender.	As máquinas podem dar respostas erradas ou mesmo não dar, mas os seres humanos também agem assim. Da mesma forma que alguns humanos são mais inteligentes que as máquinas, algumas máquinas são mais inteligentes que os humanos.

¹⁹ The popular view that scientists proceed inexorably from well-established fact to well-established fact, never being influenced by any improved conjecture, is quite mistaken. Provided it is made clear which are proved facts and which are conjectures, no harm can result. Conjectures are of great importance since they suggest useful lines of research.

4	Consciência	A máquina não tem consciência de si, portanto não tem emoções, não saberia escrever um soneto ou uma música (ou mesmo jogar xadrez).	Não é possível entrar nos pensamentos de uma máquina, da mesma forma que não é possível entrar nos pensamentos de outra pessoa e dizer o que ela sente. Seria essa uma visão solipsista, ou seja, só é possível existir a própria mente e apenas suas próprias experiências são reais, a única certeza é o próprio pensamento individual.
5	Várias deficiências	As máquinas não podem ter uma diversidade de comportamentos, como sentir gostos e saborear morangos, apaixonar-se e ter senso de humor	Turing argumenta que o mesmo acontece com o ser humano, quando, por exemplo, a mesma dificuldade de se estabelecer amizade entre homens e máquinas, como também entre brancos e negros. Que a máquina pode sim ter diversidade, desde que possua armazenamento de memória suficiente.
6	Ada Lovelace	A máquina fará tudo para a qual foi programada a fazer Ela não é criativa nem compreende algo novo.	À época de Babbage e Ada Lovelace, ainda não era possível conceber um computador eletrônico de estados discretos que pudesse ser programado com a “aprendizagem de máquina”.
7	Continuidade no Sistema Nervoso	Não se pode esperar imitar o comportamento do sistema nervoso com um sistema de estados discretos, pois ínfimas alterações na comunicação entre neurônios, por exemplo, na transmissão dos impulsos elétricos, alterariam as respostas finais.	Nas condições do jogo da imitação, o interrogador não poderá tirar proveito dessa diferença, pois a máquina pode imitar as saídas contínuas o suficiente, quando necessário.
8	Informalidade do comportamento	Não é possível criar um conjunto de regras que abranjam todas os comportamentos que possam acontecer. O ser humano pode improvisar e as máquinas não.	Ou seja, humanos não podem ser máquinas. Há uma certa confusão entre "regras de conduta" e "leis do comportamento". Turing diz ter desenvolvido um programa que fornece respostas “imprevisíveis” de números.
9	percepção extrassensorial	A ideia de que nossos corpos se movem de acordo com as leis conhecidas da física, juntamente com algumas outras ainda não descobertas, mas de certa forma semelhantes, seria uma das primeiras a cair por terra, por conta da telepatia, clarividência.	Muitas teorias científicas parecem continuar viáveis na prática, apesar de entrarem em conflito com a percepção extrassensorial; que, na verdade, é possível viver muito bem se a ignorarmos. Com a percepção extrassensorial, tudo pode acontecer. Para evitar a questão telepática no jogo da imitação, poderiam se isolar os agentes em salas especiais.

Fonte: o autor, 2026, com base no artigo de Turing, 1950

No capítulo 16, resultados e discussão, serão retomadas as objeções de Turing.

Turing retoma o argumento formulado por Ada Lovelace e introduz o conceito de “máquinas de aprendizagem”. Para o autor, a mente de uma criança poderia ser comparada a uma máquina cuja aprendizagem se inicia praticamente do zero, aproximando-se do conceito de tabula rasa. Turing exemplifica essa ideia por meio da metáfora de um caderno com folhas em branco, que vão sendo progressivamente preenchidas ao longo do processo educacional. Ainda que não se trate de uma tabula rasa

em sentido estrito, haveria, segundo o autor, a “esperança” de que exista relativamente pouco conteúdo inato no cérebro infantil.

O teste de Turing foi o marco inicial da história da IA; porém, as críticas de cunho, tanto filosófico quanto das neurociências continuam a debater sobre o que é simular processos mentais e possuir uma mente.

8. COMO PENSAM OS COMPUTADORES DE XADREZ

8.1 Teoria dos Jogos

A Teoria dos Jogos, desenvolvida por John von Neumann (o arquiteto dos computadores digitais) e o economista Oskar Morgenstern, foi a inspiração para uma das primeiras tarefas exploradas pelos cientistas da computação para os estudos da IA, o estudo dos jogos estratégicos e tomadas de decisão. O JX foi um dos laboratórios da IA, pois trata-se de um jogo determinístico observável, com informações perfeitas, regras fixas, de soma zero e que no final do jogo, os valores são sempre iguais e simétricos (Russell e Norvig, 2024). Dentre os cientistas, as primeiras iniciativas foram de Konrad Zuze e Norbert Wiener (Müller e Schaeffer, 2018; Russell e Norvig, 2024).

Estou pronto a aceitar a conjectura de Shannon de que uma máquina dessas jogaria um xadrez de alto nível amadorístico ou mesmo possivelmente magistral. Seu jogo seria assaz rígido e desinteressante, mas muito mais seguro que o de qualquer jogador humano. (Wiener, 1968, p. 172-174).

Wiener (1968), faz comentários sobre o novo tipo de máquina, o computador, “classe de máquinas assaz sinistra” (p. 172), a máquina automática de jogar xadrez.

Há algum tempo atrás, sugeri uma maneira pela qual se poderia usar a moderna máquina computadora para jogar uma partida pelo menos sofrível de xadrez. Neste trabalho, estou seguindo uma linha de pensamento que tem considerável tradição histórica atrás de si. Poe²⁰ discutiu uma máquina fraudulenta de jogar xadrez, devida a Maelzel, e a desmascarou: mostrou que era acionada por um aleijado sem pernas colocado no seu interior. Todavia, a máquina que tenho em mente é genuína e tira proveito do recente progresso das máquinas computadoras. É fácil construir uma máquina que seja meramente capaz de jogar xadrez oficial de jogar xadrez perfeito é irrealizável, pois exigiria um número muito grande de combinações. O professor John von Neumann, do Instituto de Estudos Avançados de Princeton, comentou esta dificuldade. Contudo, não é fácil, nem é irrealizável, construir uma máquina que é a mais favorável, de conformidade com algum método mais ou menos fácil de avaliação. (Wiener, 1968, p. 172-173).

8.2 O Artigo de Claude Shannon

Claude Shannon (1988), um dos pioneiros da IA, foi também o pioneiro na elaboração de ideias para os algoritmos de xadrez. O seu artigo seminal continua como inspiração para os desenvolvedores de IA e xadrez, como é o caso de Demis Hassabis e David Silver (Ensmenger, 2012; Silver *et al.*, 2018; Mitchell, 2019). Segundo Hassabis (2018), os jogos são os laboratórios para que os cientistas da computação testem seus

²⁰ Trata-se de Edgar Allan Poe, que escreveu um artigo intitulado “A Máquina de Maelzel”

algoritmos de forma rápida para que posteriormente estes algoritmos possam servir a questões do mundo real, para a saúde e a ciência.

Shannon (1988) justificou suas pesquisas com o JX para computadores afirmando:

1) O problema a ser resolvido é claramente definido, que são o xeque-mate e o movimento das peças; 2) A solução do problema não é tão trivial nem tão difícil para ser encontrada; 3) Para jogar xadrez é necessário “pensar” nas próximas jogadas, portanto, pode-se elaborar um algoritmo que as prediz ou restringir o conceito do que é “pensar”; 4) A estrutura do xadrez se encaixa na natureza digital dos computadores modernos.

Müller e Schaeffer (2018) acrescentam um quinto item, essencial para as ciências cognitivas e para a ciência da computação, que diz sobre 5) o xadrez ser um jogo em que os adversários humanos possuem uma variedade de estilos e habilidades, e que podem, portanto, avaliar a força da *engine*, e assim os pesquisadores da área terão uma introdução para artigos técnicos e científicos sobre xadrez e computação para os próximos 50 anos.

Müller e Schaeffer (2018) indicam que o artigo de 1950 de Shannon contribuiu com todos os aspectos importantes para o desenvolvimento de programas de xadrez, sobretudo para duas importantes técnicas que o preocupavam. A primeira, Tipo 1, se refere à chamada “força bruta”, ou seja, a varredura na árvore de decisão de todos os possíveis movimentos subsequentes de ambos os jogadores a ser calculada. Esta técnica tem uma limitação, o poder computacional, que reduz a profundidade de análises, o que tornaria o programa mais fraco e lento. A segunda técnica, Tipo 2, se refere à seletividade dos lances a serem escolhidos, ou seja, os “lances candidatos” (ou intuitivos) na árvore de decisão, que implica escolher quais lances devem ser selecionados para a varredura na árvore e como o algoritmo deveria saber qual lance escolher.

A primeira técnica conduz a uma exponenciação matemática de possibilidades inviável (ao menos para o poder de computação da época e sem algoritmos estatísticos e probabilísticos de lances candidatos) enquanto a segunda técnica precisava definir quais lances deveriam ser escolhidos para se caminhar pela árvore de decisão (lances intuitivos, os lances candidatos). Apresenta-se, portanto, um quebra-cabeças digno de ser usado como experimentos para a IA.

Em seu artigo, Shannon (1988) conclui que não era possível resolver o problema da mesma forma como o ser humano faz, ou seja, o computador é excelente em cálculos de variantes de xadrez e fraco em análises conceituais. Importante constatar que, dependendo do tipo de algoritmo utilizado, o da “força bruta” (Tipo 1) ou do “lance

candidato” (Tipo 2), a “personalidade” do programa podia alterar, pois um jogador “tático” (força bruta) ou um jogador de estilo “estratégico” (a seletividade na escolha de lances).

8.3 O Algoritmo Minimax

O algoritmo Minimax, cuja ideia inicial foi de Ernst Zermelo, com seu artigo publicado em 1912, é uma ferramenta relativamente simples, de fácil compreensão, rápido de ser implementado nos computadores e que serviu perfeitamente para ser implementado ao JX. O Minimax, desde os anos 1960, foi praticamente a única abordagem técnica dos programadores do JX, pois resolvia os problemas de forma satisfatória e acumulava vitórias em torneios de xadrez (Ensmenger, 2012; Russell e Norvig, 2024).

Herbert Simon (2005) comenta como é possível a construção de um algoritmo de xadrez de forma simples, visando maximizar ganhos futuros com a “boa e velha inteligência artificial” (GOFAI):

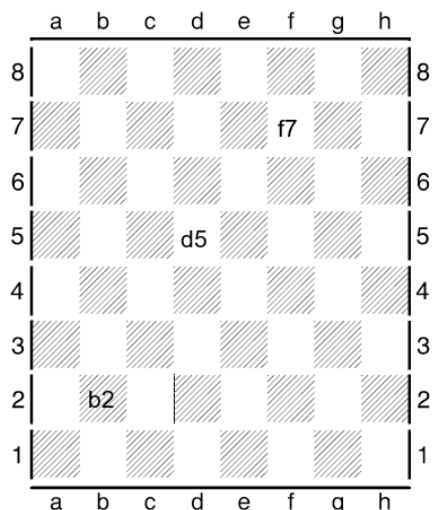
A solução do problema é uma posição futura no xadrez que maximize as chances de vitória do jogador, dado o ponto de partida. Uma árvore de movimentos legais é construída a partir da posição atual até uma certa profundidade, e cada posição alcançada é avaliada em relação à sua "qualidade" em termos de probabilidade de vitória. Por exemplo, uma função de avaliação (rudimentar) seria a razão entre o número de peças do jogador e do oponente restantes no tabuleiro. Ao minimizar retroativamente a partir das posições terminais visitadas (por exemplo, todas as posições que podem ser alcançadas por dois movimentos legais de ambos os jogadores), o "melhor" movimento pode ser selecionado. (Simon, 2005, p. 147, tradução nossa²¹).

Para se utilizar o algoritmo Minimax é necessário a construção de uma árvore de decisão e o primeiro passo é (1) a anotação e identificação do tabuleiro (Figura 11). Em seguida, (2) a identificação das peças (Figura 12). O terceiro passo (3) a execução de uma função para validar os lances legais, segundo as regras do jogo, que é o módulo gerador de jogadas e na Figura 13 é apresentado um exemplo de um Cavalo branco em “d5”, onde os “1” em cor vermelha são as possibilidades de movimento da peça. No quarto passo,

²¹ The problem solution is a future chess position that maximizes the player's winning chances, given the starting point. A tree of legal moves is grown forward from the current position to some depth, and each position reached is evaluated with respect to its "goodness" in terms of likelihood of winning. For example, one (crude) evaluation function would be the ratio of the number of player's to opponent's pieces remaining on the board. By minimaxing backwards from the terminal positions that are visited (for example, all positions that can be reached by two legal moves of both players), the "best" move can be selected.

(4) um módulo que especifica os possíveis caminhos de jogadas, ou seja, são listadas todas as possibilidades de movimentos do primeiro jogador e para cada lance é gerada uma resposta do segundo jogador, que formam a árvore a ser examinada pelo algoritmo, o Minimax. Finalmente, surge a função de avaliação que atribuirá os cálculos aos nós do fim da árvore (Ensmenger, 2012).

Figura 11 - Identificação das coordenadas

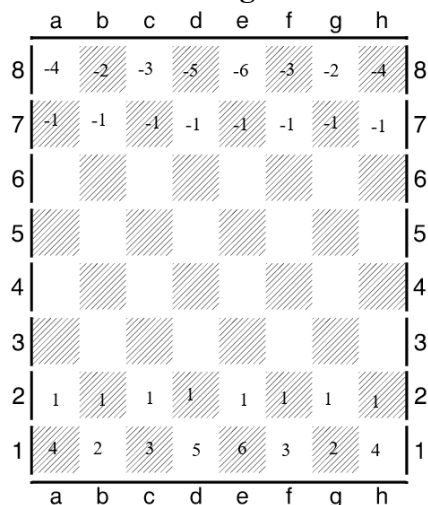


Para saber o nome das casas do tabuleiro, definidas por coordenadas, basta associar a letra do eixo das abcissas com o número do eixo das ordenadas. Por exemplo, temos as casas “b2”, “d5” e “f7”, assinaladas no tabuleiro à esquerda.

Fonte: o autor, 2026

A Figura 11 acima mostra três coordenadas em forma algébrica, que é a anotação da Federação Internacional de Xadrez. Outra possibilidade é trocar as letras das colunas por números de 1 a 8 respectivamente. Portanto, a casa “d5” seria “45”; a casa “b2” seria “22”; e a casa “f7” seria “67”.

Figura 12 - Identificação das peças



Cada número representa uma peça do xadrez.

-1 =	1 =
-2 =	2 =
-3 =	3 =
-4 =	4 =
-5 =	5 =
-6 =	6 =

Fonte: o autor, 2026

Na Figura 12 é apresentado um exemplo de como se pode representar as peças no tabuleiro. O número 6 positivo indica o Rei branco, enquanto o 6 negativo, o Rei preto.

Figura 13 - identificação do movimento do Cavalo

	a	b	c	d	e	f	g	h	
8	0	0	0	0	0	0	0	0	8
7	0	0	1	0	1	0	0	0	7
6	0	1	0	0	0	1	0	0	6
5	0	0	0		0	0	0	0	5
4	0	1	0	0	0	1	0	0	4
3	0	0	1	0	1	0	0	0	3
2	0	0	0	0	0	0	0	0	2
1	0	0	0	0	0	0	0	0	1
	a	b	c	d	e	f	g	h	

Fonte: o autor, 2026

E por fim o Minimax calculará valores para os nós acima do último grau de profundidade da árvore, sem a necessidade de novas avaliações, é apenas um procedimento mecânico de atribuir valores aos nós até o nó de onde se partiu a primeira possibilidade (o “ply”). Todos estes módulos juntos compõem o que se convencionou a chamar de “engine” (de xadrez; o motor; o sistema contendo os módulos para jogar xadrez).

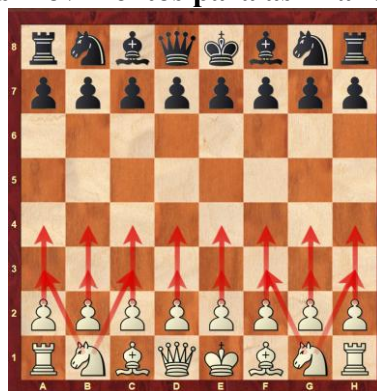
No caso do JX, a árvore de decisão se refere a todas as possibilidades de lances de um dos jogadores e as respectivas respostas do adversário. Russell e Norvig (2024) citam a complexidade sobre o tamanho da árvore de decisão: os jogadores têm, em média, 35 possibilidades de movimentação à sua escolha e uma partida tem em média 50 lances, que resulta em 35^{100} ou 10^{154} nós ou, pelo grafo, 10^{40} distintos.

Com esta informação em mente, é possível imaginar o esforço dos jogadores de xadrez durante uma única partida, pois precisam tomar decisões 50 vezes em média, e em cada decisão escolher dentre 35 possibilidades diferentes. Nos torneio de xadrez oficiais de competição, cada decisão de lance demora, em média, 3 minutos, resultando em manter a atenção por 150 minutos na partida, apenas para seus próprios lances! As partidas oficiais costumam ter duração de 4 ou 5 horas de jogo. Claro, os jogadores com alguma prática não necessitam examinar todas as possibilidades de sua própria jogada. Os

jogadores fazem uma “poda” automática para selecionar os lances candidatos. Por “poda”, se entende a exclusão de várias alternativas disponíveis para o lance a ser escolhido. Esta questão da decisão humana para o JX será explorada no próximo capítulo.

Como exemplo, a posição inicial de um JX permite a escolha dentre 20 possibilidades ou lances (Figura 14) do jogador das Brancas (pelas regras oficiais da FIDE (2023), o jogador que possui as peças Brancas inicia o jogo, que continua com movimentos alternados com o jogador das Pretas). Por sua vez, o jogador das Pretas, neste momento, possui também 20 possibilidades de respostas ao primeiro lance das Brancas (Figura 15).

Figura 14 - Possíveis movimentos para as Brancas no primeiro lance



Fonte: o autor, 2026

Figura 15 - Possíveis movimentos para as Pretas no primeiro lance



Fonte: o autor, 2026

O Minimax é um algoritmo para valorar uma sequência de jogadas em uma árvore de decisões, por exemplo de uma partida de xadrez, e assim buscar os lances mais promissores, ou seja, uma previsão dos lances a serem escolhidos para que seja possível tomar uma decisão com mais acurácia.

De início é atribuído o nome de Max (uma abreviação para valor máximo) para o melhor resultado possível em uma sequência de lances daquele que tem a vez de jogar.

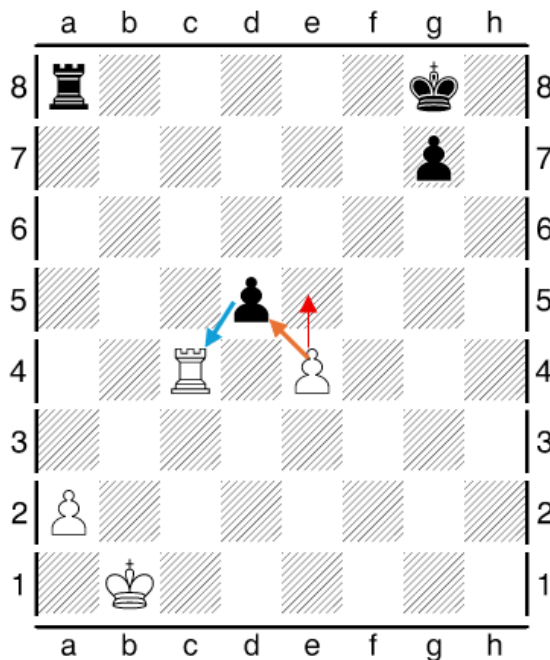
Da mesma forma, será atribuído o nome de Min (abreviação para o valor mínimo) para o melhor resultado possível em uma sequência de lances do segundo a jogar.

Ambos buscam os melhores resultados, porém, o que é bom para Min é ruim para Max e vice-versa.

Exemplo da Figura 16: se o lance que Max decide jogar ganha um Peão, o resultado será de +1 (em termos de valor material) para Max; porém, se Max decide jogar um lance que permite a perda de uma Torre, seu resultado será de -5 (em termos de valor material). Portanto, é preferível que Max decida por capturar o Peão ao invés de realizar um lance que pode perder a sua Torre. Ou seja, Max prefere capturar o Peão preto (+1; Max ganha um Peão; Min perde um Peão), mas Min prefere que Max perca sua Torre (-5; Min ganha uma Torre, Max perde uma Torre) (Müller e Schaeffer, 2018; Ensmenger, 2012).

Portanto, na Figura 16, se o Peão branco captura o Peão preto (e4xd5), então as Brancas ganham um peão (+1, Max, seta laranja). Porém, se as Brancas decidem avançar seu Peão para “e5” (seta vermelha), as Pretas podem capturar a Torre Branca de “c4” com seu Peão (d5xc4), como consequência, as Brancas perdem 5 (-5; Min, seta em azul).

Figura 16 - Opções de movimento para as Brancas



Por convenção, quando o resultado é positivo, indica vantagem na posição do jogador das peças Brancas; quando o valor é negativo, a vantagem é da posição do jogador das peças Pretas. Quanto maiores forem os valores absolutos, maior é a vantagem de um lado ou de outro (Müller e Schaeffer, 2018). Outro ponto importante é a nomenclatura que os programadores de computador dão ao movimento de cada jogador, o “*ply*” (uma redução de “*player*”, o jogador). Um lance do JX é composto de um movimento das Brancas (“*ply*”) e de um movimento das Pretas (“*ply*”) (Müller e Schaeffer, 2018; Ensmenger, 2012).

A execução do algoritmo Minimax segue de forma automática ao examinar (“varrer”) cada nó (as “folhas”) da árvore de decisão, ao atribuir valores numéricos de classificação.

Apenas os últimos nós da árvore que foi montada passam pela função, chamada “Avaliação” (será apresentado um exemplo dos parâmetros de uma “Avaliação” ao final deste capítulo), que contém as heurísticas sobre xadrez e o Minimax atribui de forma mecânica os valores intermediários aos nós, o que economiza tempo de processamento dos cálculos de avaliação em cada nível da árvore. O conhecimento a respeito do xadrez é muito bem difundido e documentado através de inúmeros livros e artigos técnicos escritos por vários mestres do jogo (Müller e Schaeffer, 2018; Ensmenger, 2012; Kasparov, 2004).

No caso do JX, as árvores de decisão são necessariamente incompletas, a não ser em posições de final de jogo com poucas peças. Por isso a função de “Avaliação” é tão importante, porém também complexa, onde a força relativa da posição tem que corresponder a um julgamento dos mais refinados, afinal, não se dispõe naquele dado momento uma conclusão definitiva sobre a posição resultante. É nesta função de “Avaliação” que muitos ajustes devem ser pensados e repensados, examinados com cuidado, para que o resultado da avaliação não vá para caminhos enganosos e sem volta.

A função “Avaliação”, portanto, atribui valores aos nós terminais da árvore construída, e em seguida, o Minimax trabalha no sentido inverso, calculando os valores dos nós intermediários acima na árvore de decisão, maximizando ou minimizando os resultados alternadamente (Müller e Schaeffer, 2018; Ensmenger, 2012; Russell e Norvig, 2024).

Cada nível da árvore representa um movimento, ou das Brancas ou das Pretas. Se o nível da árvore corresponde a um movimento das Brancas (1º “*ply*”), o algoritmo

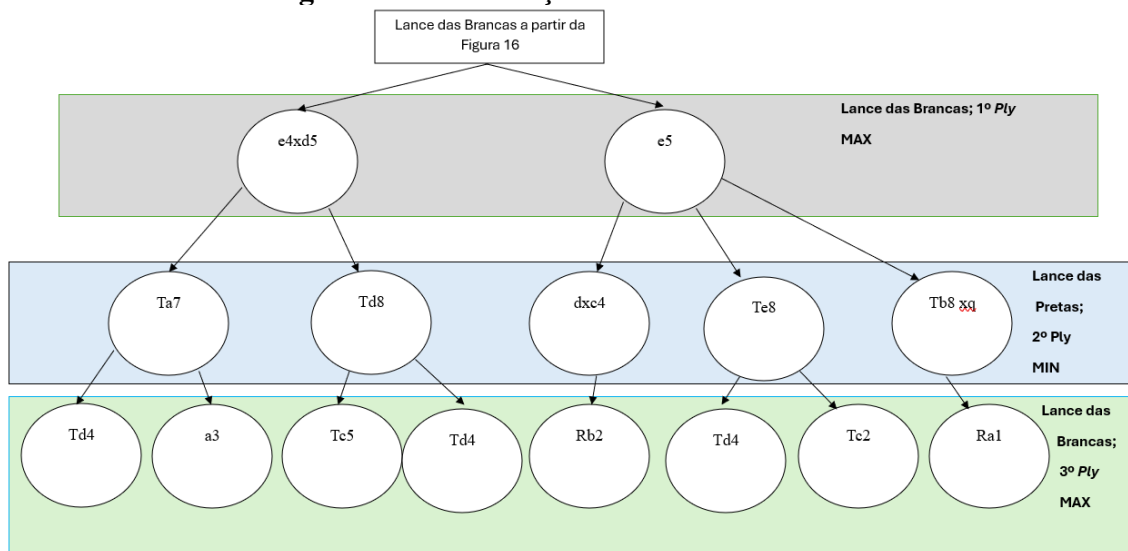
maximizará o valor de pontuação (Max) e o movimento seguinte, que caberá às Pretas ($ply + 1$), o algoritmo minimizará o valor de pontuação (Min).

Exemplo: a Figura 17 corresponde a construção da árvore de decisão (para fins didáticos, foi elaborada uma árvore com três “plys” e apenas alguns ramos da árvore visando simplificar o exemplo) a partir da posição da Figura 16. Na Figura 18 é aplicado o algoritmo de “Avaliação”. Na Figura 19 o Minimax expandirá os valores para todos os nós anteriores da árvore a partir do terceiro “ply” onde aconteceu a avaliação.

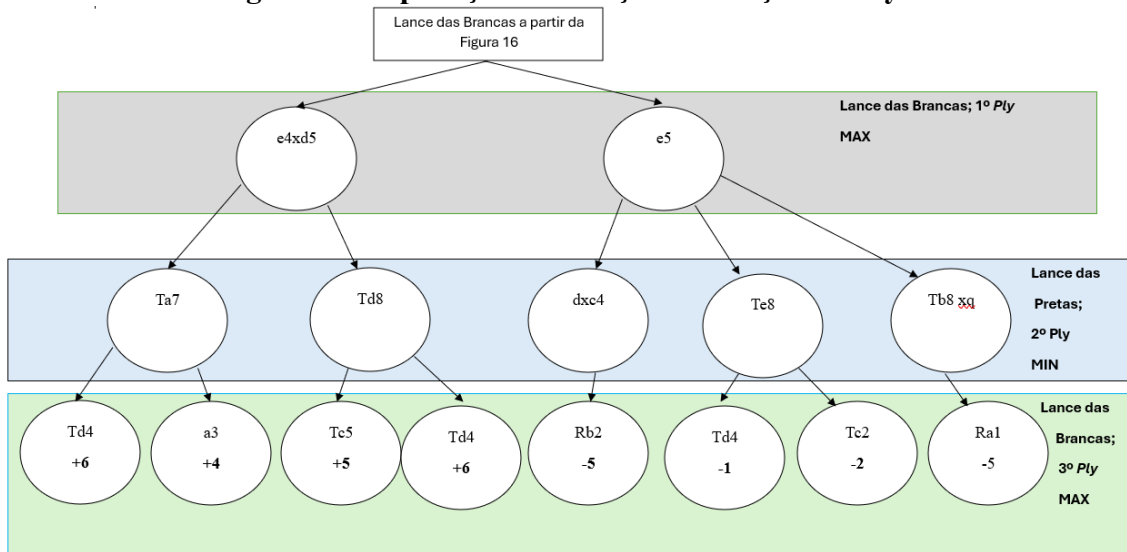
O Minimax processa os nós com valores na árvore de decisão de forma mecânica, sem a necessidade de passar pela função de Avaliação, pois os resultados são conduzidos de baixo para cima, economizando precioso tempo computacional.

Com a árvore definida, os valores numéricos atribuídos aos nós, maximizando e minimizando a cada nível, a busca pelos melhores caminhos se reduz a uma simples verificação numérica do melhor caminho. Basta percorrer os nós para achar o caminho ideal, com seu mapa pronto. Com isso, a simples melhora de processamento do computador que executa o algoritmo Minimax, produz melhor desempenho no jogo (Müller e Schaeffer, 2018; Ensmenger, 2012; Russell e Norvig, 2024).

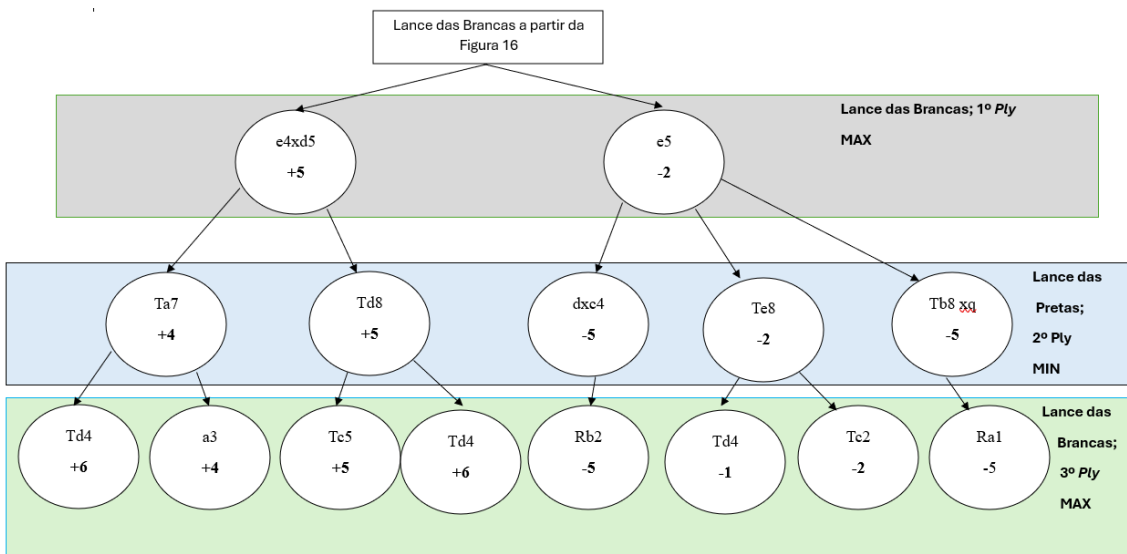
Figura 17 - Construção da árvore de decisão



Fonte: o autor, 2026

Figura 18 - Aplicação da Função Avaliação no Ply 3

Fonte: o autor, 2026

Figura 19 - Aplicação do algoritmo Minimax

Fonte: o autor, 2026

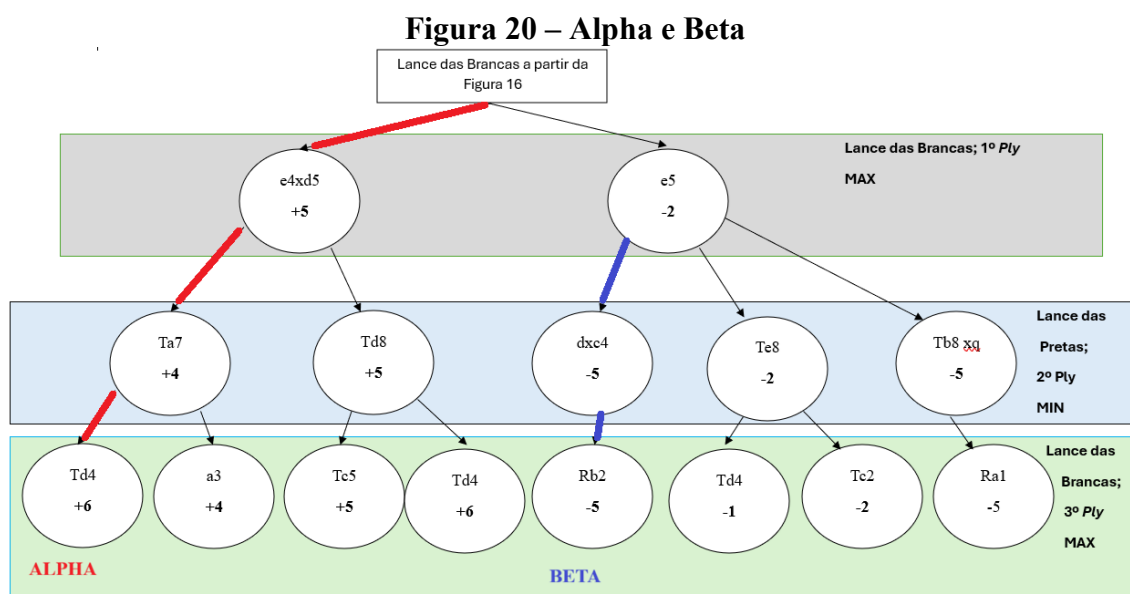
8.4 O Algoritmo Alpha-Beta Pruning

O primeiro programa de xadrez que funcionou surgiu apenas em 1957, depois de várias tentativas de experimentações em alternativa ao Minimax do Tipo 1, “força bruta”. Alex Bernstein, um forte jogador de xadrez e programador de computadores da IBM desenvolveu um programa com poda na árvore de decisão no estilo Tipo 2 de Shannon, o do “lance candidato” (intuitivo). O Minimax, porém, logo prevaleceu como o algoritmo ideal para o JX. Mas a ideia de Bernstein era de limitar o número de movimentos a serem

considerados em uma dada posição, sem alterar o resultado final, quando ainda não era feita a poda (Ensmenger, 2012).

O Minimax, por fim, ganhou uma modificação que incluiu uma melhora no algoritmo e que foi batizada de “Alpha-Beta Pruning”, devido aos esforços, em 1958, de Allan Newell, Herbert Simon e Clifford Shaw (1988). Curioso que, em paralelo, outros pesquisadores chegaram à mesma inovação, dentre os quais John McCarthy e Alan Kotok, o autor do nome Alpha-Beta, retirado de sua tese de 1962 (Müller e Schaeffer, 2018). Esta inovação no Minimax conseguiu reduzir significativamente a busca pela árvore de decisão, pois ramos inteiros da árvore eram podados, reduzindo assim os ramos da árvore a serem considerados e, por consequência, tornou possível implementar a *engine* de xadrez em praticamente qualquer computador, incluindo os primeiros microcomputadores. Houve ganhos expressivos em termos de velocidade de processamento e mais robustez de performance, desbancando de vez os algoritmos que não se relacionavam a “força bruta”, do Tipo 1 de Shannon.

A origem do nome, poda Alpha-Beta, veio de John McCarthy em sua tese de 1962 (Müller e Schaeffer, 2018) e descreve a ideia do algoritmo: Alpha é um número que representa o valor da melhor escolha que MAX pode alcançar, a avaliação do valor mais alto ao longo do caminho de MAX. Beta representa a melhor escolha que MIN pode alcançar, a de valor mais baixo ao longo do caminho para MIN (Figura 20).

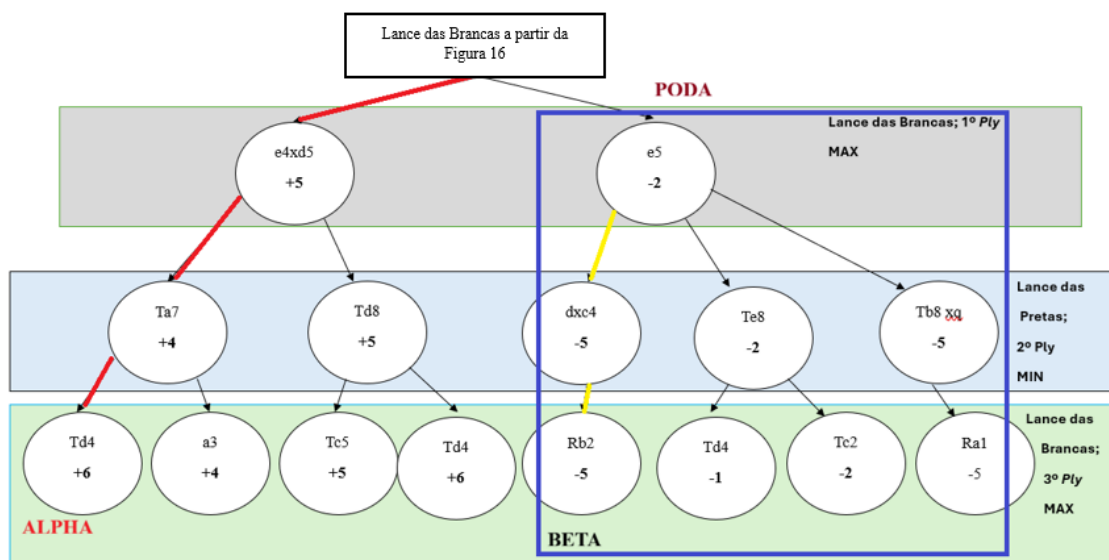


Fonte: o autor, 2026

Em outras palavras, se, por exemplo, o movimento que cabe às Brancas for avaliado de forma menos otimista que Alpha (que é o valor da melhor posição das melhores possíveis), ela é descartada, pois já se conseguiu um valor melhor que esse ramo da árvore, que é o valor de Alpha. Da mesma forma, se o movimento pessimista das Brancas é melhor que Beta (que é o pior valor da posição das mais otimistas), também esse ramo da árvore é descartado, pois as Pretas têm melhor alternativa para jogar, que é o valor de Beta. Uma ideia engenhosa, ao se comparar os valores dos diversos ramos da árvore e de se intuir a não necessidade de continuar com a profundidade de avaliação para ramos que surgem como pouco relevantes. É sempre preferível caminhar pelos melhores valores pessimistas das Brancas e os melhores valores otimistas das Pretas, ou seja, buscando os movimentos ideais de Brancas e de Pretas e descartando os intermediários.

Ao comparar estes dois números, a melhor pontuação que cada jogador pode alcançar (Alpha e Beta), é possível provar que grandes partes da árvore de busca são irrelevantes para o resultado final. Uma vez que se sabe que não se pode chegar a um resultado melhor daquele já alcançado, não precisa continuar procurando (Figura 21).

Figura 21 – Poda Alpha-Beta



Fonte: o autor, 2026

Um exemplo claro é, digamos, uma situação que contenha 10 possibilidades de movimentos. No primeiro, as Brancas alcançam o valor de +1, ou seja, o ganho de um Peão. No segundo movimento, não se consegue ganho algum. No terceiro, acontece uma

posição de xeque-mate. Pronto, não é necessário a avaliação de busca pelo restante da árvore, o melhor resultado, o xeque-mate, já foi alcançado.

Da mesma maneira, o lance das Pretas é considerado. Por exemplo, se na primeira variante da árvore, as Pretas perdem um Cavalo, na segunda perdem a sua Dama e na terceira perdem apenas um Peão, é preferível ficar com a variante que tem o melhor lance possível, neste caso para as Pretas, perder apenas um Peão.

As constantes melhorias na capacidade tecnológica de hardware e a facilidade de desenvolvimento da programação do Minimax significou em sistemas de xadrez cada vez melhores e cientificamente mensuráveis. Os mestres de xadrez com qualificações cada vez mais altas eram superados pelas *engines* com o Minimax Alpha-Beta Pruning. Os módulos relativamente independentes do Minimax eram de fácil manutenção para os desenvolvedores dos programas de xadrez e módulos como o de gerador de movimentos plausíveis, apesar da pouca inovação, ainda assim permitia melhorias importantes (Müller e Schaeffer, 2018; Ensmenger, 2012; Russell e Norvig, 2024).

O laboratório do JX foi um campo de estudos eficaz para a busca em profundidade pela árvore de decisão, as conquistas foram sendo paulatinamente acumuladas ao longo dos anos e o Minimax se mostrou uma ferramenta adequada. Porém, enquanto isso, as experiências em outras áreas da IA não conseguiram os mesmos efeitos promissores (Ensmenger, 2012).

8.5 A Função Avaliação

A implementação de heurísticas na função Avaliação exige a codificação detalhada dos fatores posicionais, conceitos estes extraídos da literatura enxadrística, que são considerados relevantes. Refinar e aumentar o número de heurísticas que são implementadas à medida que vão surgindo novos conceitos, também aumenta o tempo necessário para executar a função, diminuindo assim, o número de variantes que o programa é capaz de calcular.

A função de Avaliação se beneficia da ampla teoria de xadrez publicada a cada ano, da mesma forma que acontece com a experiência dos jogadores e a crescente base de dados com as partidas publicadas, que alimentam e sofisticam a função (Quadro 4).

Porém, ela omite detalhes sutis da experiência humana, como por exemplo, as questões culturais (Amorim, 2002; vide capítulo 14, diagramas 1 e 2).

Quadro 4 - Atributo de valores materiais das peças de xadrez

Vantagens Materiais	Valor para as Brancas	Valor para as Pretas
Xeque-Mate (o objetivo do jogo; portanto o maior ou o menor valor possível, pois o jogo acaba com sua efetivação.	+ O mais alto valor High-values	- O menor valor Low-values
Empate; posições de igualdade	Zero	Zero
Peão	+ 1	- 1
Cavalo	+ 3	- 3
Bispo	+ 3	- 3
Torre	+ 5	- 5
Dama	+ 9	- 9
Vantagem material – a diferença entre as somas das peças brancas e pretas presentes em dada posição no tabuleiro	SB – SP	SP - SB

Fonte: o autor, 2026

Amorim (2002) descreve algumas das análises estratégicas e táticas da função avaliação, que são mais difíceis de colocar em valores, que necessitam de estatísticas específicas de acordo com a posição e algumas delas estão relacionadas no Quadro 5:

Quadro 5 - Avaliação das Vantagens Estratégicas

Vantagens estratégicas:	Avaliação
Peça cravada	A peça cravada mantém o valor absoluto, o programa não analisa a cravada
Bispos de cores opostas	Avaliar este tipo de posição requer dados estatísticos e de outras peças presentes no tabuleiro, no momento em que surge para avaliação. É um problema que se insere em um contexto mais amplo.
Posições Fechadas	Planos de longo prazo, com muitas manobras de peças, a profundidade de avaliação na árvore deve ser consideravelmente maior
Transformações Posicionais	Os sacrifícios de peças, por exemplo, a troca de uma peça de maior valor por uma de menor valor, as compensações obtidas são de ordem posicional e seus efeitos concretos estão além do horizonte de cálculo.
Transição do meio jogo ao final	Poucas peças no tabuleiro, a profundidade de avaliação deve ter prioridade, as variantes a serem avaliadas são muito mais longas
Segurança do Rei	O Rei é a peça mais importante do jogo, ele deve permanecer em permanente segurança. No final do jogo ele se torna uma peça ativa

Fonte: o autor, 2026

8.6 Técnicas de Avaliação

A seguir serão listadas uma série de técnicas desenvolvidas pelos programadores de xadrez que fazem parte da função avaliação. Algumas delas são muito engenhosas, como por exemplo o “null move” e o “bit board”, que melhoraram muito os resultados da

função.

Lances Mortais – são as melhores jogadas, por exemplo, o xeque-mate; tentar um lance mortal em uma primeira busca é a chamada heurística do lance mortal (Russell e Norvig, 2024).

Ajuste na Função de Avaliação – reconhecer características da posição, como por exemplo, a estrutura de peões, a segurança do Rei, colunas abertas; Torres na sétima fila; após, atribuição de um peso a cada característica da posição as avaliações devem resultar em uma equação matemática. As dificuldades: como decidir o que é mais importante? Como avaliar as características da posição e transformá-las em números? (Müller e Schaeffer, 2018).

Tabelas de Transposição - Uma mesma posição pode ser alcançada no tabuleiro de diferentes modos. A avaliação desta posição é guardada em memória e se ela surge novamente para ser avaliada. A tabela com a avaliação é reutilizada. E se novamente ela surgir, a tabela é reutilizada mais uma vez. As transposições de posições são comuns no xadrez (Müller e Schaeffer, 2018; Russel e Norvig, 2024).

Reconhecimento de padrões – Posições padronizadas que podem acontecer no xadrez referente a diferentes jogos. Estes padrões mais comuns foram pré-programados para auxiliar na função de avaliação, uma ideia de Hans Berliner (Müller e Schaeffer, 2018).

Null Move Searching – (ou, passo o lance, que não faz parte das regras do xadrez). É uma técnica para avaliar quais são as possíveis ameaças do adversário, os piores cenários que podem acontecer na partida. Ao passar a vez, o programa recebe uma informação valiosa sobre as principais ameaças disponíveis a favor do adversário. As variantes que se mostram muito ruins não devem ser avaliadas em detrimento de outras consideradas como mais promissoras. Por exemplo, uma variante que perde uma Dama ou que leva xeque-mate é muito ruim, melhor economizar o tempo de processamento do computador utilizando-o com outras variantes mais importantes para a posição. Nem todos os ramos na árvore de decisões merecem a mesma atenção de análise, por certo, podem ser podadas ou, ao menos, não necessitam de avaliações com tanta profundidade. Economizar movimentos de busca na árvore, por exemplo, uma redução de 7 movimentos para 5 de busca na árvore, traz uma enorme economia para a *engine*. Claro que o risco da utilização desta técnica é o de ter que refazer a busca com uma profundidade maior para que seja conseguida uma melhor avaliação daquele lance ou de omitir alguma linha de busca que poderia ser importante (Müller e Schaeffer, 2018; Russell e Norvig, 2024).

Quiescência da posição - Para tomar decisões, um programa de xadrez deve conter funções que avaliem o que Claude Shannon (1988) mencionou como “posições silenciosas”, onde ocorrem, por exemplo, os xeques, que forçam respostas para salvar o Rei; as sequências de trocas de peças, pois senão o desequilíbrio material se torna decisivo; ou onde múltiplas peças estão sendo atacadas ao mesmo tempo, limitando as opções de escolha do programa. O objetivo é que deva ser feita avaliação da posição somente quando não existam os chamados lances “forçados” (Müller e Schaeffer, 2018; Russell e Norvig, 2024).

Efeito Horizonte – fixar um número de movimentos de profundidade pode causar uma má jogada por causa do limite aonde ela pode chegar com seus cálculos de variantes e deixar de fora lances que vão além do ponto de parada de avaliação (Müller e Schaeffer, 2018; Russell e Norvig, 2024).

Bit Boards – Todas as casas do tabuleiro de xadrez são transformadas em “bits”, ou seja, um bit para cada casa do tabuleiro, perfazendo um total de 64 bits. Sendo assim, a casa do tabuleiro “a1” é o bit 0; a casa do tabuleiro “a2” é o bit 1; a casa do tabuleiro “a3” é o bit 2 e assim sucessivamente até o 63º bit da casa “h8”. Uma aplicação desta ideia é aplicá-la a geração de movimentos, segurança do Rei, estrutura de peões (peões isolados, peões passados, peões atrasados, etc.). Assim, é possível calcular várias possibilidades (ou, propriedades) existentes em uma determinada posição do tabuleiro de xadrez (Müller e Schaeffer, 2018; Russell e Norvig, 2024).

Table Base (Banco de Dados ou Análise Retrógrada de finais de partida) – análise a partir de uma posição final, retrocedendo em uma sequência de lances até uma dada posição inicial. Com esta tabela, posições com sete peças ou menos, em qualquer posição (sempre com a presença dos dois Reis, Branco e Preto), já faz parte sobre como a posição inicial deve ser jogada até o xeque-mate ou empate, ou seja, um final já resolvido e armazenado em um banco de dados (Müller e Schaeffer, 2018; Russell e Norvig, 2024).

Extensões singulares – nem todas as variantes da partida devem merecer a mesma atenção. Linhas ruins devem ser menos analisadas pela *engine*. Já as linhas de análise promissoras devem merecer uma profundidade de análises maior. Ainda assim, muitos problemas não são resolvidos nos cálculos de variantes, como por exemplo, o término de um processo muito curto, desnecessário, repetição de análises, a comunicação realizada de um processador a outro, etc. (Müller e Schaeffer, 2018; Russell e Norvig, 2024).

Processadores trabalhando em paralelo – uma das melhorias nas *engines* de xadrez foi

colocar vários processadores trabalhando em paralelo para ganhar performance. A tarefa de fazer a varredura na árvore de decisão de lances de xadrez nas posições podem ser quebradas em vários computadores trabalhando em paralelo (Müller e Schaeffer, 2018; Russell e Norvig, 2024).

Extended Book – livro estendido derivado de aberturas do xadrez, que foi usado pelo *Deep Blue*, com base em mais de 700.000 partidas, das quais a *engine* da IBM tentava atrair Kasparov àquelas posições, evitando o “esforço” de busca pela árvore ainda em fase precoce da partida. Segundo Murray Campbell, um dos engenheiros de *Deep Blue*, esta técnica foi importante na segunda partida em que *Deep Blue* derrotou Kasparov, no match de 1997. (Campbell, apud Amorim, 2002).

Aprofundamento Iterativo – a ideia é o aprofundamento de uma mesma variante já analisada várias vezes, a cada vez incrementando em um nível da árvore. Motivos para que esta técnica seja usada: 1) o resultado da primeira avaliação foi irrelevante; 2) A sobra de tempo para se decidir para a jogada a ser realizada; 3) alteração na ordem do movimento prioritário a ser avaliado da árvore; 4) um resultado previamente analisado guardado na tabela de transposição. O aprofundamento Iterativo ajuda o programa a gerenciar o tempo e aumentar a eficiência da avaliação (Müller e Schaeffer, 2018; Amorim, 2002).

Função de avaliação Rápida – Em uma primeira avaliação, fatores que são fáceis de calcular (por exemplo, o material, o número de peças) e tem maior peso são processados. (Müller e Schaeffer, 2018; Russell e Norvig, 2024; Amorim, 2002).

Função de avaliação lenta - É a execução de avaliação de detalhes muito difíceis de serem avaliados e, portanto, recebem menor peso nos cálculos. Ela será dispensada de uso quando a avaliação rápida tiver uma diferença de valor que não compensa sua execução, pois seus pesos não irão influenciar na decisão sobre o lance (Müller e Schaeffer, 2018; Russell e Norvig, 2024; Amorim, 2002).

9. COMO PENSAM OS MESTRES DE XADREZ

9.1 O Investigador Pioneiro, Adrian De Groot

Os processos de pensamento dos enxadristas durante partidas de xadrez intrigam há muito os pesquisadores das ciências cognitivas. O primeiro a investigar a memória dos enxadristas foi Alfred Binet que conduziu suas pesquisas para o que poderia ser uma “memória fotográfica”, pois os mestres de xadrez jogavam partidas às cegas, sem enxergar o tabuleiro; sem dúvida uma atividade de altas habilidades de memória e abstração. Em sua investigação, Binet concluiu que os mestres de xadrez guardavam em sua memória apenas modelos abstratos do jogo e não uma “fotografia” do tabuleiro (Nicolas, 1994; Shenk, 2007; Gobet, 2019).

No JX, o processo de tomar decisões dos seres humanos envolve o reconhecimento de padrões (De Groot, 1988; Simon e Chase, 1973; Connors, Burns e Campitelli, 2011; Bratko, Hristova e Guid, 2016) e entender as consequências do que pode vir a acontecer em uma determinada posição.

Adrian de Groot foi o pioneiro a elaborar um método de investigação para saber o que pensavam os mestres de xadrez em suas partidas. Seu trabalho de doutorado foi referência para outros pesquisadores, dentre eles um dos pioneiros da IA, Herbert Simon (Gobet, 2019).

Para realizar seus experimentos, De Groot utilizou questionários e pediu para que os jogadores falassem em voz alta enquanto jogavam, contando o que estavam pensando durante a partida. De Groot trabalhou com vários grandes mestres de xadrez e identificou que eles, principalmente, buscavam reconhecer padrões advindos de estruturas que foram analisadas e compreendidas em seus estudos do jogo e de suas experiências; com isso os mestres desenvolveram a capacidade de compreender com rapidez a essência da posição, mesmo observando-a por poucos segundos (Gobet, 2019).

O principal interesse de De Groot era o processo de tomada de decisão na escolha dos lances do jogo. A sua hipótese era que os mestres de xadrez consideravam sequências mais longas de jogadas (mais profundidade na árvore de decisão) comparando com amadores (Gobet, 2019).

Suas expectativas foram frustradas: os mestres ainda conseguiam jogar movimentos melhores, porém as possibilidades de movimentos considerados e a profundidade na busca não eram diferentes dos amadores de xadrez. Os melhores jogadores tiveram vantagem em identificar rapidamente as soluções mais promissoras da

posição, o que restringiu a quantidade dos “lances candidatos” a serem avaliados. Enquanto isso, os amadores demoraram muito mais para achar os lances candidatos. Enquanto o jogador amador demorava 15 minutos em suas avaliações, o campeão mundial demorava apenas 5 segundos! O que diferenciou mestres e amadores, concluiu De Groot, foi a percepção da posição, o fator essencial para a perícia no jogo (Gobet, 2019).

O segundo experimento realizado por De Groot ficou famoso: ele apresentava uma determinada posição de xadrez no tabuleiro aos jogadores e após alguns segundos a retirava, pedindo que os jogadores remontassem a posição. Os mestres de xadrez conseguiam reconstruir a posição com facilidade, praticamente com 100% de acerto. Quanto aos amadores, dependia de seu nível de jogo. Quanto mais fraco o jogador, mais errava na reconstrução da posição. Havia uma hierarquia na capacidade de lembrança da posição, do grande mestre ao amador, as escalas de acerto iam diminuindo conforme a força enxadrística (Gobet, 2019).

Aspectos perceptivos se entrelaçavam com os aspectos dinâmicos das peças no tabuleiro, ou seja, não era meramente o estado estático das peças, não eram as peças individualmente, mas o grupo de peças, a compreensão das possibilidades dinâmicas do jogo que os mestres de xadrez percebiam. A busca visual dos mestres pelo tabuleiro se fixava com frequência na intersecção das casas onde atuavam as peças (Gobet, 2019).

Esta é uma das principais contribuições de De Groot, ele demonstrou que a percepção, a memória e a resolução de problemas estão intimamente interligadas. Percepção e cognição não tem uma fronteira clara (Gobet, 2019).

Outro resultado importante do trabalho de De Groot foi notar que, tanto os mestres quanto os jogadores amadores investigavam um dos primeiros movimentos da situação problema várias vezes, que denominou como “aprofundamento progressivo”. Ou seja, os jogadores analisavam um dos principais movimentos várias vezes, intercalando com outros ou não, mas retornando as análises daquele mesmo movimento (Gobet, 2019).

Reinvestigar o movimento base inicial, é uma tentativa de aprofundar e reavaliar com mais precisão e compreensão daquela jogada. Os enxadristas reexaminam os movimentos ciclicamente, em um processo que envolve observação, teste, avaliação. Segundo Fernand Gobet (2019) não se trata de um processo ineficiente e repetitivo, mas sim adaptativo, pois ajuda a memória de curto prazo, que é limitada, podendo estas variantes serem armazenadas na memória de longo prazo, o que torna mais fácil a análise da posição com o auxílio de novos recursos cognitivos.

De Groot traz exemplo da pesquisa científica sobre o aprofundamento progressivo e da maneira como os humanos lidam com problemas complexos e tomam decisões: o pesquisador científico explora uma solução possível, examina outra, retorna à primeira solução e assim por diante. Um processo que pode ser de apenas alguns minutos ou de anos (Gobet, 2019).

Ainda a respeito do aprofundamento progressivo, surgem diferenças nas habilidades dos mestres e amadores. Quando os jogadores reavaliam determinado ramo da árvore de busca, esta pode ser realizada de duas formas, uma de forma imediata e outra não imediata (intercalada por outras variantes). Os mestres utilizam uma estratégia do “ganhar-ficar” e “perder-mudar”, termos da Teoria dos Jogos. Se a reavaliação for positiva, vale a pena rever a sequência dos movimentos daquele ramo da árvore de forma mais aprofundada, porém, se a reavaliação for negativa, é porque a linha é desfavorável e a abandonam (Gobet, 2019).

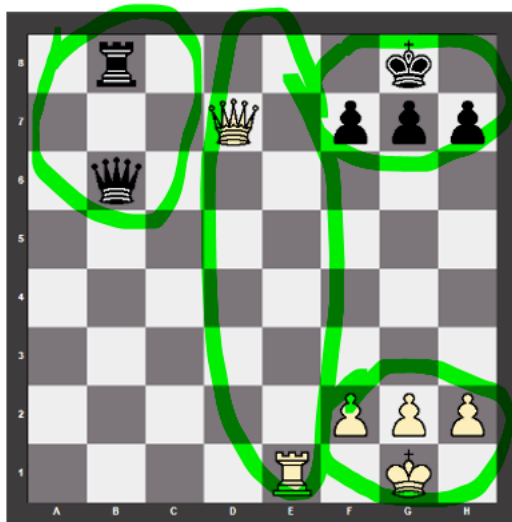
De Groot divide o pensamento dos enxadristas em quatro fases, que ele nomeia de orientação, exploração, investigação e comprovação. A fase de orientação é aquela onde são colhidas informações sobre o problema a ser resolvido com uma primeira breve avaliação, uma preliminar; na exploração os jogadores reduzem o número de possibilidades consideradas críticas para duas, após analisarem algumas variantes; na investigação, estes dois movimentos selecionados são analisados de forma mais profunda, onde os jogadores buscam uma confirmação para suas preferências, ou seja, desejam que as linhas que estão investigando estejam certas, em uma tentativa de que um dos movimentos é melhor que o outro; na fase de comprovação, os jogadores recapitulam suas conclusões e verificam se estão corretas (Gobet, 2019).

9.2 Chunks

Simon e Chase (1973), com base nas pesquisas de De Groot e o reconhecimento de padrões, investigaram e desenvolveram uma nova teoria, a dos “*chunks*”, pequenas porções do tabuleiro que, ao se juntarem, formam conceitos. Eles descobriram que os mestres de xadrez não reconheciam padrões de todo o tabuleiro, mas sim de pequenas porções que se agrupavam para formar a totalidade do jogo no tabuleiro. Eles combinaram métodos experimentais, modelagem computacional e ideias da IA para realizar a série de experimentos que culminou em seus artigos de 1973 e que se tornaram fonte de inspiração de pesquisas nas décadas seguintes (Gobet, 2019).

A teoria dos “*chunks*” pressupõe que os jogadores de xadrez codificam seus conhecimentos na sua memória de longo prazo em “porções” (os *chunks*, blocos), as unidades perceptuais, que podem ser tratadas como um todo, quando recuperadas da memória (Figura 22). Os primeiros blocos são pequenos, mas blocos maiores são construídos incrementalmente usando os blocos menores. Para um novato no xadrez, apenas as peças individuais formam os blocos e à medida que se vai estudando e aprendendo, obtendo conhecimento, estes blocos vão aumentando, crescem em agrupamento de peças (Gobet, 2019).

Figura 22 - Chunks



Fonte: o autor, 2026

A busca altamente seletiva dos mestres é fruto destes *chunks* armazenados, depois de muita aprendizagem. Os *chunks* são o recurso para a construção de planos, táticas, técnicas e até mesmo a rapidez com a qual os exímios jogadores encontram os melhores lances no jogo, identificando padrões, os melhores lances e as possíveis sequências de jogadas (Gobet, 2019).

Para concluir, o conhecimento adquirido pelos mestres de xadrez fica armazenado na sua memória de longo prazo como fragmentos perceptivos que estão vinculados a possíveis ações e a busca altamente seletiva (Gobet, 2019).

9.3 A importância de Conhecer o Pensamento do Mestre de Xadrez

Desenvolvida na década de 1950 por Herbert Simon, a racionalidade limitada propõe que os humanos tomem decisões que são boas o suficiente, mas não ótimas, devido aos limites impostos por seus recursos cognitivos (Russell e Norvig, 202). A racionalidade limitada faz parte do arsenal técnico para tomar decisões dos jogadores de xadrez.

Chabris (2017) observa a importância de se estudar o comportamento humano através de jogos como o xadrez, afirmando que os jogos são como laboratórios experimentais da psicologia. Chabris afirma que mesmo os grandes mestres de xadrez cometem erros, enquanto as *engines* conseguem achar todos os erros humanos. Os computadores que jogam xadrez hoje em dia são parceiros críticos para entender as tomadas de decisão humanas.

Chabris (2017) comenta em sua tese de doutorado, que entrevistou um grande mestre de xadrez, que lhe revelou seus pensamentos durante uma partida, sobre conceituar posições ao invés de fazer apenas cálculos de variantes. O mestre de xadrez consome seu tempo fazendo análises, estratégias, relações lógicas entre as peças. Chabris reforça a ideia de que os inúmeros livros de xadrez ensinam conceitos que devem ser aprendidos pelos enxadristas, como por exemplo, sobre buscar a iniciativa no jogo, o controle de espaço, peças atacadas ou mal posicionadas e como fazer para melhorá-las. Questões estas sobre abstrações que são a base do pensamento dos grandes jogadores.

Portanto, Chabris (2017) propõe que se deva elaborar um protocolo com perguntas a grandes jogadores para saber o que eles pensam durante a partida, se são cálculos de variantes ou se são análises de conceitos abstratos, ou seja, entender o que se passa no tabuleiro para tomarem suas decisões de lances a serem executados durante uma partida de xadrez.

9.4 Como Pensa Gary Kasparov

Gary Kasparov (2007) fez uma reflexão à pergunta comum, “o que se passava na sua cabeça?” (Kasparov, 2007, p. 2), a respeito de como tomava decisões no xadrez.

A reflexão que Kasparov fez foi de como suas decisões o estimulavam a evoluir para o crescimento pessoal, a autoavaliação, o autoconhecimento, sobre se desafiar e aos outros e assim aprender a tomar as melhores decisões. Kasparov não se referia apenas ao xadrez, mas como o xadrez o ajudou na própria vida pessoal.

Kasparov tece comentários a respeito das influências psicológicas individuais e do caráter emocional expresso em suas decisões; de como o xadrez é um instrumento para examinar essas influências, pois os mestres são forçados a rever as decisões tomadas nas partidas e refletir como chegaram até elas (Kasparov, 2007),

Kasparov também menciona que a prática em tomar decisões melhora o processo intuitivo e inconsciente. Emite sua opinião sobre a importância de fazer planos, que o jogador visualiza uma posição futura de 10, 20 lances a frente, estabelece objetivos e só depois elabora os lances passo a passo e é só nesta fase que surgem os cálculos de variantes (Kasparov, 2007).

Sobre a criatividade, Kasparov afirma que ela é complementar à lógica e ao raciocínio dedutivo; é necessário muito estudo (conhecimento adquirido), prática, memória e imaginação, ou seja, “pensar fora da caixa”. A imaginação, segundo Kasparov, é sinônimo de fantasia, de fugir dos padrões normais, onde o cálculo de variantes vem apenas após as ideias imaginadas (Kasparov, 2007).

O reconhecimento de padrões é mencionado por Kasparov (2007) da seguinte forma:

Quando tentamos resolver um problema, nunca começamos do zero. Instintivamente, até de maneira inconsciente, procuramos um paralelo no passado. Calculamos a autenticidade dos paralelos, e verificamos se podemos produzir uma receita semelhante com esses ingredientes um tanto diferentes. (KASPAROV, 2007, p. 71).

Sobre os cálculos que os mestres fazem e a busca pela árvore de decisão durante a partida, Kasparov comenta:

Quando contemplo meu lance, não começo percorrendo imediatamente a árvore de decisão. Primeiro, tenho de considerar todos os elementos na posição para poder definir uma estratégia e desenvolver objetivos intermediários. Preciso ter todos esses fatores em mente quando finalmente começar a calcular as variantes, para saber que resultados são favoráveis. Experiência e intuição podem orientar esse processo, mas uma rigorosa base de cálculos continua se mostrando necessária. (...) Minha experiência me orienta a escolher dois ou três lances candidatos como objetivo. Em geral, um deles é descartado rapidamente como lance inferior, e outro é considerado para tomar seu lugar. Em seguida, começo a expandir a árvore com um lance por vez, analisando as reações prováveis e meus lances de resposta. (Kasparov, 2007, p. 59-60).

Kasparov acrescenta ainda que, para ser eficiente, além da disciplina mental de descartar os lances que são insatisfatórios, a árvore de decisão deve ser podada com frequência e saber quando parar na profundidade.

Kasparov enfatiza a avaliação que faz das posições de xadrez, dizendo que o lance é apenas o resultado concreto da equação que foi desenvolvida e entendida. A determinação dos fatores relevantes, avaliá-los e estabelecer o melhor equilíbrio possível entre eles.

O primeiro elemento que Kasparov indica para a avaliação é a relação de material, a quantidade de peças presentes no tabuleiro, mas sempre observando a segurança do Rei, pois se ele receber um xeque-mate, de nada adianta ter muito material a mais. O Rei é o essencial da posição.

O segundo elemento destacado por Kasparov é o tempo, não aquele que mede a contagem de horas, mas sim o “número de passos necessários para alcançar um objetivo” (Kasparov, 2007, p. 97).

Claro que também existe o tempo do relógio de xadrez, que limita os jogadores a realizar suas partidas dentro do limite estabelecido pelo torneio. O relógio é um fator psicológico do jogo, um motivo de pressão sobre os jogadores, que muitas vezes apressa as decisões, em especial quando o tempo vai se esvaindo (Kasparov, 2007).

Para entender a questão do tempo no JX, é necessário entender a dinâmica do jogo. Cada jogador tem a vez de executar o seu lance. Se fosse possível realizar, por exemplo, três lances seguidos, a vantagem do jogador seria enorme, pois este poderia posicionar as suas peças de uma forma que não deixaria muitas opções de defesa para o adversário. O jogador que conduz as peças Brancas tem a vantagem do lance inicial, do tempo a mais. Portanto, as Brancas realizam uma ação e as pretas reagem à ação (Kasparov, 2007).

O fator tempo no xadrez tem um caráter dinâmico; um lance equivocado de um dos lados pode significar a perda de tempos valiosos para movimentar as peças para as melhores casas. Em muitas situações, um jogador pode “sacrificar” um de seus peões para simplesmente poder ter um tempo a mais, manejar suas peças de forma a trocar o material para poder ocupar uma casa desejada. Kasparov alerta sobre este tipo de avaliação, a troca de material por tempo (Kasparov, 2007).

O próximo elemento considerado por Kasparov (2007) é denominado qualidade, os fatores de longo prazo versus fatores dinâmicos da posição. Isto indica que, a posição das peças no tabuleiro, por exemplo, um cavalo centralizado ao invés de um cavalo na borda, é mais dinâmico, age mais rapidamente, quando necessário, próximo da ação do jogo. Um cavalo muito afastado do centro das ações perde muito tempo em chegar a tempo para colaborar com as demais peças de seu exército. Uma peça aprisionada, sem

mobilidade, continua sendo uma peça, porém seu valor relativo é muito menor. Kasparov chama isso de “conceito do nunca”, uma possível dificuldade de avaliação dos algoritmos de xadrez.

O material só tem o valor de sua utilização. O tempo para a ação só é importante se nos ajudar a aumentar a utilidade de nosso material. Ter uma hora de sobra todos os dias seria agradável para quase todo mundo, mas não para um homem numa cela de prisão. Devemos utilizar o tempo para aperfeiçoar nosso material, não apenas para adquirir mais. Material em si e tempo desperdiçado são igualmente inúteis para atingir nossos objetivos (Kasparov, 2007 p. 104).

Sobre as decisões no JX, Kasparov comenta: “Para executar uma estratégia, é necessário tomar decisões. Para avaliações se transformarem em resultados, elas devem gerar decisões. Depois das fases de preparação, planejamento, análise, cálculo e avaliação, temos de escolher uma linha de ação” (Kasparov, 2007, p. 167).

Kasparov questiona sobre como examinar o processo de decisão e de como ajustá-lo. A importância da qualidade da informação é ponderada, pois ao selecionar os lances a serem examinados, que são as informações percebidas pelo mestre de xadrez, pode-se deixar de fora outros lances.

Porém, aumentar o leque de possibilidades tem um custo, um luxo limitado pelo fator do tempo do relógio que controla as decisões durante a partida. Cinco ou seis lances lógicos são possíveis em determinadas posições, embora haja normalmente dois ou três.

Portanto, a primeira tarefa a ser realizada é limitar a amplitude da busca na árvore. Para o enxadrista experiente, o cálculo preliminar de variantes facilita o trabalho. Só depois de constatar que estas primeiras opções não satisfazem é que outros lances podem ser considerados. Essa mudança gera perda de tempo, além de ser uma escolha psicológica difícil de ser realizada, pois as suposições iniciais não foram satisfatórias e não há garantias de que as novas possibilidades possam ser melhores daquelas já examinadas.

Resulta disso dois tipos de padrões opostos e destrutivos, quais sejam, escolher por um caminho já avaliado e bem conhecido ou fazer uma escolha precipitada de uma opção nova e desconhecida. Conclui que é melhor errar no lado conhecido a “cair às cegas no abismo na esperança de pousar na nuvem” (Kasparov, 2007, p. 172).

Por essa razão é tão importante começar por pelo menos duas opções e tempo suficiente para avaliar ambas. Mergulhar em uma análise profunda de apenas uma das alternativas pode deixá-lo sem tempo bastante para examinar outras, e você ficará preso entre esses dois padrões negativos. Quando tiver terminado, já é muito tarde, pois o tempo terá acabado (Kasparov, 2007, p. 172).

Sobre os lances candidatos e poda da árvore de decisões, Kasparov tem sugestões, conforme o momento do jogo, mas afirma ser essencial limitar o número de lances candidatos e podar a cada passo dado para não tornar inviável o aprofundamento da análise o suficiente para conseguir alguma solução plausível. “Examinar cinco opções diferentes para dois lances à frente não é melhor nem pior que examinar apenas duas opções para cinco lances, dependendo do problema e da posição que tivermos em mãos” (Kasparov, 2007, p. 174).

Enfim, o tempo do relógio é o fator para adotar diferentes posturas e se decidir por qual estratégia utilizar para a avaliação e quantidade dos lances candidatos, mas sempre com uma poda da árvore de decisão rigorosa para não se perder em longas variantes que podem ser inúteis no fim das contas.

Por último, é importante lembrar do processo de iniciativa, ou seja, ter a vantagem de comandar as ações. Ter a iniciativa significa poder ampliar as possibilidades dos lances candidatos, em outras palavras, ter mais opções de escolha na árvore de decisão, enquanto os lances adversários são mais restritos. Para aquele que possui a iniciativa, a vantagem de ter o lance de tempo a mais e o seu uso significa ter mais flexibilidade e mobilidade, desvios (Kasparov, 2007).

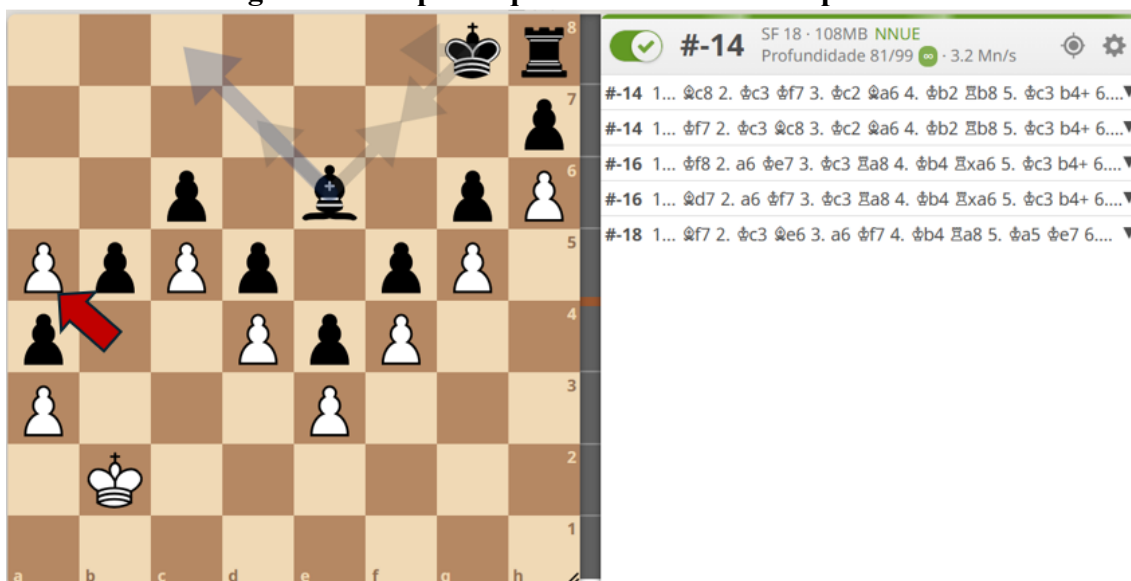
10. O LABORATÓRIO JOGO DE XADREZ

O JX é um laboratório de pesquisas amplo, investigado em disciplinas como por exemplo, sociologia, etnologia, filosofia, matemática e neurociências. Porém, o xadrez é utilizado com mais frequência e eficácia nos campos da IA e na psicologia. Na IA, o xadrez é muito utilizado no aprendizado de máquina e nos algoritmos de busca e na psicologia em temas sobre percepção, memória, aprendizagem, tomada de decisão (Gobet, 2019). Não à toa, é considerado a drosófila da psicologia cognitiva.

Serão apresentados aqui alguns exemplos que apontam as diferenças nas decisões entre humanos e *engines*. Nestes exemplos surgem as controvérsias sobre as decisões da IA especialista em xadrez, onde a solução do problema não é resolvida com o puro cálculo de variantes. Os seres humanos, ao contrário, percebem o que deve ser feito, contextualizam e solucionam os problemas que exigem imaginação e não cálculos, probabilidades e estatísticas. Este é o ponto fraco das IAs de xadrez e o ponto forte dos seres humanos.

O primeiro exemplo provém dos estudos de Roger Penrose (2021). Em seu livro, “Sombras da Mente”, de 1994, Penrose examinou uma posição de xadrez que fora construída para demonstrar que compreender a posição do jogo era mais importante que a vantagem material. Penrose inseriu a posição no computador da época a fim de examinar a solução que lhe seria fornecida. O resultado foi que a *engine*, “*Deep Thought*”, antecessor de *Deep Blue*, cometeu um erro, capturou uma torre que o fez perder a partida. *Deep Thought*, escolheu o lance do ganho de material para diminuir o prejuízo da posição (Figura 23). A escolha que *Deep Thought* deveria ter feito era simplesmente mover o Rei, pois, apesar da enorme desvantagem material, as pretas estavam sem lances para fazer desbloquear a posição.

Figura 23 - Captura que foi um erro do computador



Fonte: sítio www.lichess.com, construído pelo autor, 2026

Resultado: o lance levou as brancas à derrota, pois a variante calculada por *Deep Thought* não chegava ao ponto onde ficava claro o erro em tomar a torre das pretas. A conclusão lógica de Penrose foi que, simplesmente, *Deep Thought* estava apenas calculando variantes e, portanto, não conseguia entender as sutilezas da posição, ao contrário de qualquer mestre de xadrez.

Neste estado do tabuleiro, as peças pretas têm uma vantagem enorme, pois ainda possuem duas torres e um bispo. Porém, é fácil para o jogador das peças brancas evitar a derrota, simplesmente movendo seu rei do seu lado do tabuleiro. (...) Como pode um jogador de xadrez maravilhosamente eficaz fazer um movimento tão estúpido? A resposta é que o *Deep Thought*, além de ter sido municiado com uma quantidade considerável de “conhecimento dos livros”, foi programado somente para calcular movimento – com uma profundidade razoável – e tentar melhorar sua situação no tabuleiro. Em nenhum momento possuía qualquer entendimento real do que uma barreira de peões poderia alcançar – nem, de fato, poderia ter qualquer entendimento genuíno seja do que for que faça”. (Penrose, 2021, p. 81).

Os computadores da época em que Penrose fez este teste estavam muito longe do poder de processamento de *engines* atuais. Para atualizar tais conclusões, foi utilizada a *engine* “Stockfish 18 NNUE”, no sítio da web livre (www.lichess.com) que permite executar esta versão do algoritmo de xadrez, considerado um dos mais fortes da atualidade, que é uma combinação do Minimax com redes neurais. O resultado é apresentado (Figura 24), a seguir:

Figura 24 - A primeira escolha do Stockfish 18 não é a captura



Fonte: sítio www.lichess.com, construído pelos autores, 2026

Desta vez, a *engine* “Stockfish 18” tomou uma decisão diferente: apesar de apresentar cálculos com uma pontuação negativa para as brancas (-4.3), a *engine* não escolheu o lance que perde a partida. Ela continua sem entender a posição do jogo, seus cálculos ainda deixam a *engine* como se estivesse perdendo a partida de Brancas. Houve sim uma melhora significativa de seus cálculos, mas isto não quer dizer que ela esteja em posição de saber que a partida está empatada.

No exemplo seguinte (Figura 25), a *engine* “Stockfish 18” atribuiu cálculos para os quatro possíveis lances do rei branco como -10,9 (cada ponto corresponde ao equivalente a um peão a mais ou a menos; neste caso, é como se as brancas tivessem o equivalente a uma força de 10,9 peões a menos, ou uma Dama a menos, cada Dama vale 9 pontos). O quinto lance possível das brancas, o avanço de seu peão para a casa “c7”, a *engine* calcula um xeque-mate para as pretas em nove lances. “Stockfish 18” não cai na armadilha de mover seus peões, mas seus cálculos mostram valores muito negativos às brancas, ou seja, a consideração de uma derrota. “Stockfish 18” vai continuar jogando sem saber que a partida está empatada, ao contrário de um jogador humano, que rapidamente percebe a situação.

Figura 25 - Composição de Penrose, 2017



Fonte: sítio www.lichess.com, construído pelos autores, 2026

O próximo exemplo vem de uma posição de partida que repercutiu contra a força das *engines*, advinda do confronto entre a *engine* “Rybka”, com as peças brancas e o grande mestre estadunidense, Hikaru Nakamura, em 2008 (chess.com, 2021). Nakamura jogou de forma a esperar o erro da *engine*, pois “entendia” ser uma posição de empate.

Segundo as regras da Federação Internacional de Xadrez (FIDE, 2023), quando não é realizado lances no tabuleiro com movimentos ou capturas de peão por pelo menos cinquenta lances seguidos, a posição é considerada empate²². A *engine* não conseguiu “entender” o que acontecia, e por ter o “conceito” em seus algoritmos de que possuía material a mais, cujo cálculo era a seu favor o suficiente para não deixar que a partida empatasse, capturou a torre preta e acabou perdendo a partida.

Para rever as análises, mais uma vez, foi utilizada a *engine* “Stockfish 18” e o resultado se alterou, apresentando mais uma vez o avanço das IA de xadrez. O “Stockfish 18” não caiu na armadilha de Nakamura, e seus cálculos apontam uma leve desvantagem (-0.5) contra, e sua primeira escolha. Depois de alguns minutos de processamento, deixou o lance perdedor além das cinco primeiras opções (Figura 26).

²² As condições estabelecidas pela FIDE para que o jogo seja declarado empatado, segundo os artigos 9.3 e 9.6 do regulamento: 50 lances, empate opcional, se algum dos oponentes solicitar; 75 lances, empate declarado pelo árbitro.

Figura 26 - Rybka versus Nakamura, 2008



Fonte: sítio www.lichess.com, construído pelos autores, 2026

Para um mestre de xadrez humano, a posição é de empate. Apesar da evolução das *engines*, para o Stockfish 18, entender a posição ainda está muito distante da avaliação conceitual feita pelos humanos. Nas posições de teste mostradas, os computadores mais fortes têm performance inferior a jogadores humanos de capacidade muito modesta.

O exemplo seguinte é do campeonato mundial de xadrez de computadores de 1977. A posição da Figura 27 aconteceu após o 34º lance da *engine* Duchess.

Com o Rei preto em xeque, o lance óbvio, que salta à vista é o lance 34...Rg7. Porém, a *engine* Kaissa, jogou o surpreendente 34....Te8, entregando a Torre, ao invés de fugir com o Rei, ou seja, um “bug”, uma falha enorme na programação de Kaissa. Os especialistas em xadrez presentes à sala onde se realizava a partida, deram boas risadas, devido ao lance incompreensível da *engine*.

Dentre os presentes, que também riu bastante, estava o ex-campeão mundial Mikhail Botvinnik, engenheiro que trabalhou com os computadores soviéticos de xadrez, considerado o “pai” da escola soviética de xadrez e um dos professores de Gary Kasparov.

Os programadores, ao verificarem o porquê de um lance tão ruim, ficaram surpresos, pois o lance de Kaissa adiava o xeque-mate da *engine* Duchess. Nenhum especialista na sala, nem mesmo Botvinnik, percebeu que o lance Rg7 era refutado pelo brilhante lance 35.Df8+ com xeque-mate em mais quatro lances para as Brancas.

Ficou evidenciado o quão poderosas eram as *engines* com a chamada “força bruta” do poder computacional (Müller e Schaeffer, 2018).

Figura 27 - Duchess versus Kaissa, 1977



Fonte: sítio www.lichess.com, construído pelos autores, 2026

11. DEEP BLUE E KASPAROV

Quando Kasparov foi derrotado por *Deep Blue*, Douglas Hofstadter (2005), se decepcionou e admitiu a “força bruta” da máquina e lembra que a vida é diferente do xadrez:

Para repetir, a analogia é tentadora — mas a vida é muito diferente do xadrez. A vida não tem regras fixas, nenhum objetivo único abrangente, nenhuma grade cristalina na qual se mover, nenhuma torre, cavalo e peão rigidamente definidos. Muito pelo contrário, a vida consiste em milhares de esperanças, medos e sonhos que interagem simultaneamente e de forma complexa, cada um deles envolto em sua própria nuvem única e intangível de confusão intelectual e emocional. (Hofstadter, 2005, p.194, tradução nossa²³).

A narrativa aqui se utiliza dos vários conceitos propostos pelos autores da TAR (Callon, 1986; Callon, 1999; Law, 2007; Tonelli, 2016; Latour, 2011; Latour, 2012; Latour, 2013). Nela é possível percorrer os fios condutores que ligam atores humanos e não humanos, cujas ações mobilizam a dinâmica dos fatos, efeitos de translação (ou tradução) que jogam luz a aspectos da ciência e da tecnologia em construção. As controvérsias que surgem da narrativa “homem versus máquina”, ao se percorrer os caminhos por fios que ligam atores humanos e não humanos foram obtidas através de pesquisas em jornais da época, filmes sobre o assunto e em artigos científicos sobre IA e o jogo de xadrez, sobre o embate entre Kasparov e *Deep Blue*, que são recursos para narrar este estudo de caso (Latour, 1997).

É a partir de um laboratório inusitado, o jogo de xadrez, que acontecem a ciência em ação. Em um palco onde se situa a mesa de jogo com tabuleiro, peças e o relógio que controla o tempo das jogadas de ambos os competidores. Em dois lados opostos à mesa, duas cadeiras, uma correspondente ao campeão mundial de xadrez da época e outra ao técnico da IBM que simplesmente executa os lances escolhidos pela máquina, exibidos por um terminal de computador que fica ao seu lado.

As consequências sociais em razão da nova ciência e tecnologia, a IA e a arquitetura dos 256 processadores do computador a calcular em paralelo duzentos milhões de movimentos do jogo por segundo (Bory, 2019), devem ser exploradas de forma simétrica e imparcial, ou seja, conhecer os acertos conquistados pelos lances efetuados na

²³ To repeat, the analogy is tempting—but life is very different from chess. Life has no fixed rules, no single overarching goal, no crystalline grid in which to move, no sharply constrained rooks, knights, and pawns; quite the contrary, life consists in thousands of simultaneously and intricately interacting hopes, fears, and dreams, every last one of them surrounded by its own unique, intangible cloud of intellectual and emotional blur.

partida de xadrez por *Deep Blue*, mas também as falhas que ficaram escondidas até os reveladores depoimentos da equipe que o construiu, quase vinte anos após o confronto.

Segundo Latour (2000), o conceito de “caixa preta” é o da tecnologia já estabelecida, portanto já são de domínio público os aspectos de sua construção. *Deep Blue*, ao contrário, revelou apenas uma arquitetura inovadora, que foi cobiçada por indústrias, como a farmacêutica, que viram nele uma tecnologia com dimensões econômicas e financeiras de grande potencial, mas guardou segredos de seus algoritmos, sua “alma”.

“*Deep Blue*” foi uma “caixa preta”; os programas de computador que atuavam em seus circuitos integrados não foram revelados, nem aos cientistas da informática externos à equipe da IBM, muito menos aos assessores de Kasparov. Por isso, os movimentos no jogo de xadrez efetuados por *Deep Blue* exerceram uma função tão importante quanto às do humano Kasparov.

Computadores, mídias impressas e eletrônicas, mercado financeiro, peças e tabuleiros de xadrez, relógios eletrônicos, internet (um ator recém-nascido à época) e humanos, atuaram de forma social em uma rede complexa, repleta de fatos não tão óbvios, que repercutiram em vários segmentos sociais, econômicos e culturais.

A política não ficou tão evidente, como na “Guerra Fria”, no enfrentamento de xadrez entre o americano Bobby Fischer e o russo Boris Spassky em 1972, àquela época, a supremacia do poder intelectual, capitalistas ou comunistas, mas ainda sim Kasparov era o russo a ser derrotado pela indústria americana. “*Deep Blue*”, nesta narrativa, é um computador, perfeitamente entendido pela sociedade, como uma máquina de alta capacidade de processamento de dados. Porém, em seu interior é executado um “*software*”, um sistema de algoritmos, que foram guardados em segredo pela IBM.

A IBM, desde o início de suas atividades, marcou presença no campo político-militar. O campo científico da computação iniciou com a forte influência de militares em seus projetos durante a segunda guerra mundial, fosse para decifrar os códigos secretos dos inimigos, como também para realizar os cálculos da bomba atômica. A IBM foi uma participante ativa na então promissora indústria de computadores (Edwards, 1995).

No relato que segue, as críticas feitas à TAR por Sismondo (2010), são abordadas de forma que, ao menos para a IA e computadores, a TAR é bem adequada às investigações do tema.

O confronto entre o então campeão mundial de xadrez, Gary Kasparov e o computador “*Deep Blue*”, da corporação estadunidense IBM, em 1997, ficou conhecido como a disputa “Homem versus Máquina”, que teve uma intensa divulgação na mídia, dentre as quais se destacaram jornais, revistas, jornais televisivos e a então recém-criada Internet (que chegou ao Brasil em fins de 1995). Para se ter uma ideia da intensidade de aficionados e curiosos que acompanhavam as partidas de xadrez entre o humano e o não humano, houve uma quebra de recordes de acessos nos sítios da IBM e das páginas oficiais do evento na internet. A IBM, interessada na ampla divulgação do evento, não poupou esforços no sentido de instigar revistas e outras publicações midiáticas (Jayanti, 2003; Bory, 2019) a divulgarem o evento.

Esta disputa entre *Deep Blue* e Kasparov foi marcada por inúmeras controvérsias. Kasparov ficou insatisfeito com a IBM que não disponibilizou os algoritmos do computador para que fossem analisados, acusou lances realizados por *Deep Blue*, dizendo que não eram de computadores, mas sim de humanos (como na segunda e sexta partidas jogadas entre ambos) além de outras questões sobre o confronto (Kasparov, 2017; Bory, 2019).

Os trabalhos dos engenheiros, cientistas e consultores do supercomputador *Deep Blue* não foram simplesmente “descobertos” de como deveria ser a arquitetura inovadora da máquina. O trabalho em paralelo dos 256 processadores de *Deep Blue* veio pela necessidade em conseguir um processamento de dados excepcional para vencer o campeão Kasparov (Bory, 2019). Tratou-se, portanto, de uma construção social, não só da arquitetura do computador, mas também dos aperfeiçoamentos de um outro ator importante da rede, o algoritmo conhecido como “Minimax” já existente à época, que em sua busca pelos lances das variantes da partida na árvore de decisão, podava as ramificações que eram consideradas menos importantes para a escolha do movimento de *Deep Blue*, tornando-o mais eficiente em sua busca (Ensmenger, 2012).

As práticas e culturas dos cientistas da computação para a construção de *Deep Blue* tiveram que ser submetidas às exigências da empresa IBM. O actante IBM fez com que os cientistas trabalhassem com as práticas impostas em busca de vencer Kasparov, de forma materialista e relacional, o contexto foi determinado pela empresa e os objetivos eram claros, derrotar o campeão.

Naquela época, a famosa indústria de computadores IBM estava sentindo os efeitos de concorrentes da informática que traziam como foco os computadores pessoais,

ao contrário da IBM, que esteve presente nas maiores empresas do mundo, fornecendo seus computadores de grande porte (*mainframes*) (Bory, 2019). Sendo assim, atores como a Microsoft e a Apple exerceram grande pressão, de forma indireta, ao confronto entre Kasparov e *Deep Blue*.

A IBM procurava algum fato que pudesse fazer com que ela voltasse a se destacar no cenário internacional de informática. E foi o que realmente aconteceu.

Com a vitória de *Deep Blue* ao final do “*match*”²⁴ de seis partidas contra Kasparov, as ações da IBM tiveram um aumento de valor vertiginoso, elevado a patamares como há dez anos a empresa não conseguia.

Outro detalhe importante daquele momento foi a perspectiva de a IBM entrar em um novo campo estratégico, o da IA. A tecnologia de IA estava estagnada naquela época, e a vitória de *Deep Blue* reascendeu com toda a força as pesquisas neste campo revolucionário da informática.

A quebra de paradigma, a IA superando o humano no JX, teve um grande efeito social, uma nova ameaça ao trabalho das pessoas, pelo temor de serem substituídas pelas máquinas dotadas de IA e da provável queda da disponibilidade de empregos no mercado de trabalho (Bory, 2019).

O protagonista desta narrativa, que foi colocado na posição de herói e que, como ele próprio se considerou, defendia a supremacia da IH sobre a máquina, era Kasparov. O campeão de xadrez chegou a pedir que fosse colocada uma bandeira ao seu lado (ao invés da russa) que representasse a humanidade contra o vilão da história, *Deep Blue*. Porém, outros atores ganharam protagonismo, dentre eles o próprio *Deep Blue*, no decorrer do confronto.

Por exemplo, logo na primeira partida (do total de seis), *Deep Blue* fez um lance que deixou Kasparov perplexo. A máquina mostrou uma intencionalidade improvável. A jogada foi tão enigmática, que, apesar da vitória de Kasparov nesta partida, apenas dois movimentos depois do lance improvável de *Deep Blue*, deixou o campeão a pensar que estava diante de um novo tipo de inteligência, muito maior que a sua e que fazia cálculos de variantes com uma antecipação de jogadas (profundidade) além do que ele podia fazer e que não podia ser compreendida. Kasparov ficou perturbado pelas próximas cinco partidas que ainda seriam disputadas.

²⁴ Conjunto de partidas do enfrentamento, que no caso foram seis.

Vinte anos depois, em um documentário de Frank Marshall (2014) sobre o confronto, da rede de televisão ESPN, o depoimento de um outro actante da rede, o grande mestre de xadrez, o norte americano Joel Benjamin elucidou o mistério daquele lance realizado por *Deep Blue*: simplesmente, houve um erro no algoritmo (um erro humano dos programadores), que deixou o programa em “*looping*”²⁵ e havia uma estratégia de que, caso acontecesse algum erro, o algoritmo deveria escolher aleatoriamente um lance possível de ser realizado para que efetuasse um lance qualquer de sua árvore de decisão e sair do *looping*.

A segunda partida então, foi dramática. *Deep Blue* passou a dividir o protagonismo com Kasparov. Os analistas que seguiam os lances da partida, ficaram encantados com a estratégia demonstrada pela máquina. Próximo ao final da partida, *Deep Blue* realizou “um lance humano”, segundo Kasparov, abdicando do cálculo material (e realista) para realizar um lance “intuitivo”, de longo prazo (profundidade estratégica), cujo valor era difícil de ser avaliado.

Kasparov, em situação complexa no jogo, perdeu a partida. Na entrevista dada logo após, Kasparov disse ter percebido “a mão de deus”²⁶ nas jogadas de *Deep Blue* (Kasparov, 2017). Kasparov estava convencido que grandes mestres auxiliavam com seus próprios lances as decisões de *Deep Blue*.

Um fato curioso efluuiu de expectadores da internet, horas depois: se Kasparov tivesse jogado mais alguns lances, ao invés de abandonar (termo enxadrístico onde se reconhece a vitória do adversário) a partida por se considerar em posição perdida, ele, com lances precisos, poderia ter empatado.

Outras análises posteriores mostraram que *Deep Blue* cometeu um erro de cálculo no fim da partida, improvável para uma máquina que calculava 200 milhões de lances por segundo, e, ao invés do lance ganhador, executou um lance inexplicável para uma máquina. *Deep Blue* empatou o “*match*” em 1 a 1, mas ainda restavam quatro partidas para que o confronto se definisse (Marshall, 2014; Kasparov, 2017).

As terceira, quarta e quinta partidas mostraram Kasparov cabisbaixo e jogando fora de seu estilo agressivo e o resultado foi que elas terminaram empatadas.

²⁵ Circunstância em que as instruções do programa de computador se repetem sem sair de uma condição determinada, executando instruções em círculo sem as finalizar.

²⁶ Uma alusão às palavras ditas pelo jogador de futebol argentino Maradona, em jogo da copa do mundo de futebol de 1986, quando fez um gol irregular, com a mão.

A última partida, a sexta e decisiva, *Deep Blue* novamente protagonizou o espetáculo e fez um lance que não era o primeiro da sua lista de escolhas: um sacrifício de peça²⁷, bem conhecido daquela abertura, até por amadores de xadrez, mas que não é a preferência dos cálculos materialistas do computador.

Kasparov caiu numa armadilha muito conhecida, talvez não acreditando que *Deep Blue* conduzisse o jogo por aquela variante.

Joel Benjamin, mais uma vez, jogou luz a mais esta controvérsia: ele havia programado algumas respostas automáticas evitando as escolhas do computador, prevendo a possibilidade de Kasparov seguir por estes caminhos.

Outro integrante da equipe da IBM, o grande mestre de xadrez Miguel Illescas declarou em uma entrevista para a revista britânica especializada em xadrez, “*British Chess Magazine*”, que Benjamin teria inserido o fatídico lance nas bases de dados de *Deep Blue*, na manhã do dia da partida, o que seria ter acertado na loteria, tão raro seria saber a escolha de Kasparov para os lances iniciais da última e decisiva partida do “match”.

Mais uma vez, sem entender o que havia ocorrido, Kasparov saiu esbravejando, muito nervoso, acusando a empresa IBM e a equipe de programadores, engenheiros e construtores de *Deep Blue* de terem trapaceado (Marshall, 2014; Kasparov, 2017).

Um dos truques de *Deep Blue* foi que ele havia sido programado para se envolver em jogos psicológicos. Uma artimanha dos técnicos da IBM para confundir Kasparov. Quando havia um lance óbvio, *Deep Blue* aguardava um tempo antes de realizar seu lance, mesmo já tendo decidido qual iria jogar. Por outro lado, lances ótimos de Kasparov, que exigiam mais processamento, eram resolvidos de forma mais rápida. Um truque psicológico embutido nos algoritmos da máquina.

O grande mestre Miguel Illescas da equipe de *Deep Blue*, comentou: “essa estratégia tem um impacto psicológico, pois a máquina torna-se imprevisível, o que era o objetivo principal”. (Eysenck e Eysenck, 2023, p. 22).

Uma inscrição importante deste laboratório, as “logs”²⁸ do processamento de dados de *Deep Blue* não foram liberadas pela IBM, portanto não se pode saber o que efetivamente aconteceu nos cálculos de variantes processados pelo algoritmo e as controvérsias ainda hoje permanecem em aberto. Por outro lado, as partidas realizadas

²⁷ Entrega deliberada de uma de suas peças sem compensação material, ou seja, sem obter peça de igual valor do adversário.

²⁸ Relatório com os passos executados pelos algoritmos.

entre Kasparov e *Deep Blue* perduram como fonte histórica para consultas e análises, uma inscrição duradoura advinda do laboratório “partida de xadrez”. Estas inscrições são provas do que ocorreu entre os actantes que protagonizaram o embate desta rede sociotécnica. Muitos livros já foram escritos, inclusive pelo próprio Kasparov (2017), cheios de análises, observações e inúmeras discussões, e as controvérsias permanecem.

Nesta rede apresentam-se vários actantes, dentre os quais os dois protagonistas, Kasparov e *Deep Blue*, mas outros actantes fundamentais para a narrativa, a IBM, Joel Benjamin, Miguel Illescas, o engenheiro e arquiteto de sistemas Feng-Hsiung Hsu, os demais cientistas da equipe de computação, a bolsa de valores, a internet, o público presente ao *World Trade Center* em Nova Iorque, local da disputa, e a mídia em geral que evidenciou o confronto “Humano versus Máquina” no topo das notícias ao redor do mundo.

Kasparov, após o confronto, jogou um dos melhores torneios de sua carreira e depois de se aposentar dos tabuleiros em 2005 passou a se dedicar a seus livros de xadrez e sobre estratégia e planejamento comparando o jogo com a administração de uma empresa. Se candidatou à presidência da Rússia e foi até preso pela polícia de Vladimir Putin. Kasparov também prestou serviços como palestrante por diversas vezes para funcionários da corporação IBM, inclusive no Brasil, sobre IA, planejamento e administração (Kasparov, 2007; Kasparov, 2017).

Deep Blue, ao contrário, apesar da fama de ter conseguido ser o primeiro computador a derrotar um campeão mundial de xadrez, mudando assim o paradigma sobre IH e IA, ao menos no JX, foi desmontado logo em seguida ao confronto. Nos dias de hoje, sua tecnologia ficou totalmente ultrapassada, tanto em relação à sua arquitetura quanto a lógica fria e materialista de seus algoritmos e poderia ter se tornado uma peça de museu.

Quando as duas opiniões se enfrentam no tabuleiro, a criatividade de um indivíduo extremamente talentoso competirá contra o trabalho de gerações de matemáticos, cientistas da computação e engenheiros. Acreditamos que o resultado não revelará se as máquinas podem pensar, mas sim se o esforço humano coletivo pode superar as maiores realizações dos seres humanos mais talentosos – Trecho de uma declaração de Feng-Hsiung Hsu e os outros membros da equipe do Deep Thought/*Deep Blue*. (Kasparov, 2007, p. 242).

Utilizar a metodologia TAR para descrever o evento que ficou conhecido como a disputa “Homem versus Máquina”, permite ampliar o cenário que impactou os avanços das pesquisas com inteligência artificial, em especial na quebra de paradigma da

supremacia humana sobre as máquinas, um evento simbólico disputado no tabuleiro de xadrez.

A vitória da tecnologia sobre o campeão mundial em um campo considerado de exclusividade da inteligência humana, o xadrez, propiciou a quebra de paradigma. A inteligência humana fora superada pela inteligência artificial em um jogo considerado de altas habilidades cognitivas para os participantes.

Deep Blue se mostrou como uma tecnologia encomendada pela multinacional IBM, a fim de que a corporação pudesse se reestabelecer como a maior empresa de computadores mundial. Portanto, *Deep Blue* foi uma construção social, confirmando o materialismo relacional, um dos princípios da TAR. As propriedades reais e intrínsecas de *Deep Blue* se mostraram relativas, aperfeiçoamentos ocorreram entre as partidas jogadas no confronto, longe dos olhos de Kasparov e sua equipe.

As práticas e culturas dos cientistas e engenheiros tiveram que se adaptar às exigências da empresa, quase que eliminando a subjetividade de seus trabalhos. Os fatores econômicos contextuais impuseram a maneira de conduzir os trabalhos.

A narrativa da rede sociotécnica encontrou e mudou o holofote direcionado aos vários actantes, ao tornar ativo tanto o protagonismo do humano quanto o de *Deep Blue*, de forma simétrica e imparcial. O provável herói, o campeão Kasparov, saiu ofuscado pela nova tecnologia. Outros actantes ganharam importância na narrativa para se compreender todo um contexto que cercava o confronto “Homem versus Máquina”, como por exemplo o grande mestre de xadrez Joel Benjamin, a bolsa de valores, a internet e é claro, a IBM.

Nesta narrativa não houve acordo entre os cientistas, as disputas pela nova tecnologia derivada da inteligência artificial foram ferozes e abarrotadas de controvérsias. Empresas surgiram alguns anos depois para desenvolver tecnologias de IA mais modernas como o “Google” e a inglesa “*Deep Mind*”, baseadas em redes neurais, evidente colaboração da interdisciplinaridade, tendo as neurociências e a filosofia como disciplinas colaboradoras principais.

Os atores desta rede sociotécnica surgiram em diversos pontos: o grande mestre Joel Benjamin já havia disputado partidas com Kasparov e havia uma disputa subjetiva em segundo plano.

O próprio Kasparov demonstrou um conflito em especial, ambivalente, pois mudou seu estilo, costumeiramente agressivo, apesar de sua primeira vitória na primeira

partida do “*match*”. Kasparov sabia que poderia ser um evento que marcaria a história, a primeira vez que um campeão mundial de xadrez seria derrotado por uma máquina construída para jogar xadrez.

Uma controvérsia possivelmente escondida, teria Kasparov usufruído de benesses financeiras para perder o confronto, ou o seu enorme ego não permitiria tal barganha?

A TAR se mostrou uma ferramenta eficiente para analisar as controvérsias de um dos mais famosos eventos midiáticos no final do século XX.

12. A REVOLUÇÃO DO ALPHAZERO

A empresa DeepMind, fundada por Demis Hassabis (mestre de xadrez), desenvolveu em 2016 o programa AlphaGo, criado para disputar partidas de Go, o mais antigo dos jogos conhecido pela humanidade e cuja complexidade é significativamente maior do que a do xadrez (Murray, 1910).

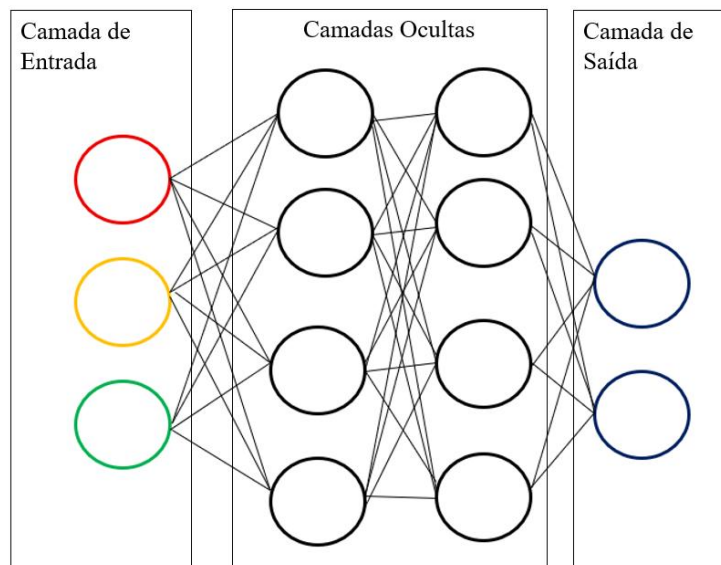
O AlphaGo alcançou a façanha de vencer alguns dos mais fortes jogadores de Go da época, entre eles Lee Sedol, considerado uma lenda do Go tanto quanto Kasparov é para o xadrez. O êxito do AlphaGo levou os pesquisadores a conceber diferença uma nova *engine*, o AlphaZero, capaz de aprender e jogar, não apenas Go, mas também xadrez e shogi (uma modalidade japonesa de xadrez) (Mitchell, 2019).

O AlphaZero utiliza técnicas de redes neurais convolucionais (*convolutional neural networks* – ConvNet), combinadas com aprendizado por reforço (*deep reinforcement learning* - DRL), e o algoritmo de busca em árvore de Monte Carlo (*Monte Carlo Tree Search* – MCTS) (Mitchell, 2019).

As ConvNets são compostas por unidades de entrada (*inputs*), que correspondem aos dados que alimentam a rede; por camadas ocultas, formadas por nós (ou neurônios) interconectados entre si; e por unidades de saída (*outputs*), que representam os resultados do processamento dos dados de entrada pelas camadas internas (Mitchell, 2019).

Os neurônios da rede estão conectados e recebem pesos, que determinam a influência de um neurônio sobre os demais nas camadas subsequentes. Quanto maior o número de camadas ocultas, mais profunda é considerada a rede neural (Mitchell, 2019).

O chamado aprendizado profundo, DL, refere-se justamente à capacidade de ajustar iterativamente esses pesos ao longo de múltiplas camadas, de modo a otimizar o desempenho da rede neural (Mitchell, 2019). A Figura 28 mostra uma rede neural com duas camadas.

Figura 28 - Rede Neural Convolutucional

Fonte: o autor, 2026

Esse ajuste de pesos ocorre por meio de processos de aprendizado baseados em *feedback*, nos quais os erros ou acertos das saídas geram sinais de correção. No caso do aprendizado por reforço, tais sinais assumem a forma de recompensas, que orientam a atualização dos pesos, que são atualizadas por meio de algoritmos de retropropagação, permitindo à rede melhorar progressivamente sua performance (Mitchell, 2019).

O aprendizado por reforço é um método cuja origem remonta aos trabalhos de Arthur Samuel, engenheiro da IBM desde 1949, que desenvolveu um dos primeiros programas capazes de aprender a jogar Damas. O sistema de Samuel incorporava mecanismos hoje associados ao aprendizado por reforço, combinados com a busca em árvore do tipo *Minimax*, o uso de uma função de avaliação e estratégias de autojogo, por meio das quais o programa aprimorava seu desempenho a partir da experiência acumulada (Mitchell, 2019; Russell e Norvig, 2024).

A ideia de aprendizagem por reforço tem raízes na psicologia behaviorista, cuja doutrina enfatiza o reforço de comportamentos considerados adequados e a punição de comportamentos inadequados (Mitchell, 2019).

No contexto das RNA, quando o desempenho do agente é avaliado como positivo, o sistema recebe sinais de recompensa que orientam o ajuste dos pesos sinápticos; quando o desempenho é insatisfatório, os sinais recebidos conduzem a ajustes que reduzem a probabilidade de repetição daquele comportamento (Mitchell, 2019).

A aprendizagem por reforço caracteriza-se como uma estratégia flexível, na qual não é necessário que humanos ensinem explicitamente regras ou enumerem todas as circunstâncias possíveis de um ambiente (Silver *et al.*, 2018; Mitchell, 2019).

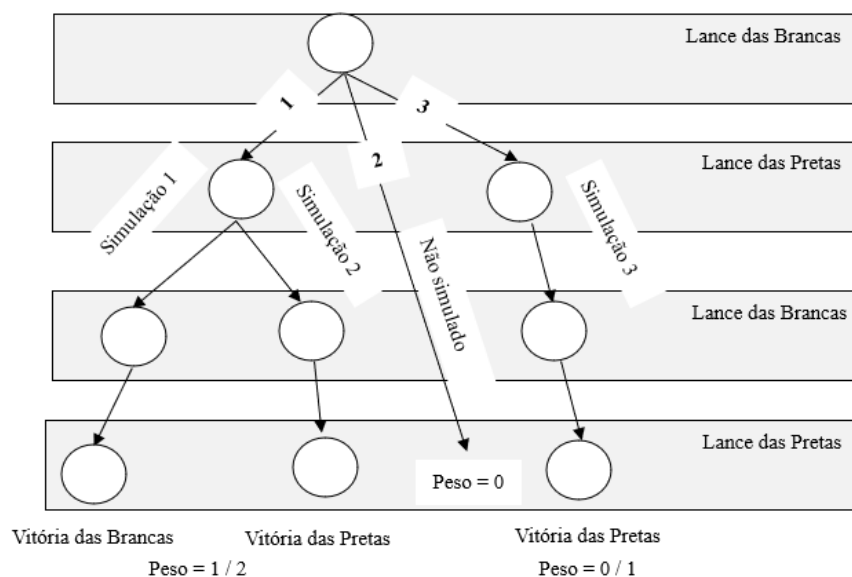
Diferente da aprendizagem supervisionada, esse método não requer conjuntos de dados rotulados por especialistas humanos, baseando-se, em vez disso, na interação contínua do agente com o ambiente e na avaliação de suas ações por meio de recompensas (Mitchell, 2019).

O método da busca em árvore de Monte Carlo é um algoritmo baseado em amostragem aleatória, cujo nome remete aos métodos probabilísticos desenvolvidos a partir de analogias com jogos de azar, como a roleta dos cassinos (Mitchell, 2019).

Seu funcionamento consiste em explorar o espaço de decisões por meio de simulações estocásticas, nas quais diferentes sequências de ações são avaliadas e estatisticamente agregadas. A partir de um número suficiente de simulações, o algoritmo é capaz de estimar, com boa aproximação, os valores esperados das ações disponíveis, orientando a escolha das jogadas mais promissoras (Russell e Norvig, 2024). A Figura 29 apresenta um exemplo da MCTS.

Diferente de abordagens baseadas em redes bayesianas explícitas, o MCTS não requer o conhecimento prévio da distribuição verdadeira de probabilidades, apoiando-se, em vez disso, na convergência estatística obtida pelo grande número de amostragens realizadas durante a busca (Russell e Norvig, 2024).

Figura 29 - Busca na Árvore de Monte Carlo



Fonte: o autor, 2026

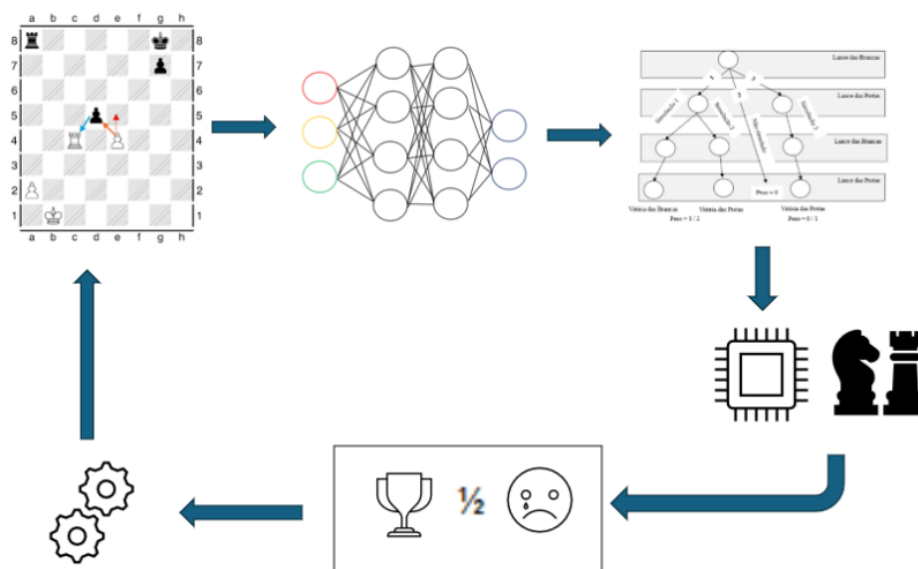
A ConvNet gera valores como indicador de quais movimentos devem ser escolhidos para a “roleta” do MCTS, o que restringe o número de lances candidatos a serem jogados, aumentando as probabilidades de sucesso (Mitchell, 2019).

Além do método de busca em árvore de Monte Carlo, o AlphaZero utiliza redes ConvNets, treinadas para estimar valores aproximados associados às diferentes opções de jogadas possíveis a partir de uma determinada posição, servindo como guia inicial para a expansão da árvore de Monte Carlo (Mitchell, 2019).

A ConvNet produz, de um lado, uma distribuição de probabilidades que indica quais movimentos são mais promissores e, de outro, uma estimativa de valor da posição, funcionando como um mecanismo de priorização semelhante a uma “roleta” probabilística no processo de seleção das jogadas. Esse direcionamento reduz significativamente o espaço de busca, restringindo o número de lances candidatos considerados e aumentando a eficiência e as chances de sucesso do algoritmo (Mitchell, 2019).

Em essência, o AlphaZero aprende jogando milhões de partidas contra si mesmo, processo no qual o sistema é instanciado simultaneamente nos dois lados do tabuleiro, Brancas e Pretas (Figura 30).

Figura 30 – Arquitetura do AlphaZero



Fonte: o autor, 2026

A posição do JX é informada para a ConvNet; a ConvNet fornece as probabilidades dos lances com os seus respectivos valores para a Busca na Árvore de Monte Carlo; a MCTS simula algumas variantes de forma aleatória e faz uma escolha

final; o resultado pode ter sido uma vitória, empate ou derrota; AlphaZero aprende com o resultado; pesos serão atualizados para retropropagação (reforço). Portanto, AlphaZero não calcula todas as possíveis variantes, mas sim apenas aquelas escolhidas ao acaso e aprende com o resultado.

Esse auto jogo constitui a fase central de treinamento do algoritmo. A partir das posições geradas durante as partidas, cada instância do agente utiliza o método de busca em árvore de Monte Carlo para construir e explorar múltiplas sequências possíveis de jogadas futuras, associando a cada uma delas estimativas probabilísticas de sucesso ou vitória (Mitchell, 2019).

Com base nessas estimativas, o algoritmo seleciona as sequências mais promissoras, aquelas com maior valor esperado, e orienta suas escolhas de lances de acordo com essa avaliação.

O mesmo procedimento ocorre de forma simétrica para o outro lado do jogo, que também busca maximizar suas probabilidades de êxito. Como resultado, diferentes sequências de jogadas são comparadas e ponderadas por sua “força”, isto é, pelos valores atribuídos pela rede neural e pela busca em árvore. Esse processo se repete continuamente ao longo de milhares de iterações, permitindo que o sistema aprenda e refine progressivamente estratégias mais eficazes.

David Silver *et al.* (2018) afirmou que esse método é particularmente eficaz porque, no processo de autojogo, o oponente apresenta sempre um nível de dificuldade adequado ao estágio de aprendizado do próprio sistema.

O AlphaZero tem um estilo muito intrigante e intuitivo no JX, sacrificando material em troca de compensações posicionais de longo prazo que a maioria dos humanos não consegue ver. Assim como o que aconteceu com o Go, o AlphaZero parece mostrar que a compreensão humana do xadrez, apesar de séculos de prática e estudo, é bastante limitada (Sadler e Regan, 2019).

O AlphaZero não apenas executa jogadas eficientes, mas revela um modo distinto de tomada de decisão, baseado na avaliação probabilística de futuros possíveis e na otimização contínua de valor. Essa forma de decidir, embora formalmente matemática, apresenta analogias estruturais em relação aos processos decisórios humanos, especialmente no que diz respeito à intuição, à antecipação e ao julgamento sob incerteza (Sadler e Regan, 2019).

O Stockfish NNUE é uma mistura dos algoritmos tradicionais como o Minimax, mas que contém RNA. O Stockfish é considerado como a melhor *engine* atual.

Já a *engine* LeelaZero, ou LC0, é construída apenas por RNA, estilo AlphaZero.

Embates entre o Stockfish e o LC0 mostra uma diferença de estilos muito interessante, e que são utilizados pelos mestres de xadrez para suas preparações de aberturas (os lances iniciais do xadrez) em todas as competições das quais eles participam (Sadler e Regan, 2019; Ribeiro-Doggers, 2025).

13. REVISÃO BIBLIOGRÁFICA

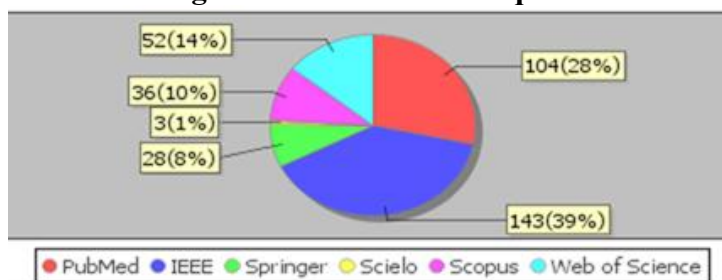
13.1 Resultados da Revisão

A pesquisa por palavras-chave obedeceu ao seguinte argumento genérico de busca (em inglês): (“Chess” or “game of Chess” or “chess game”) and “decision making” and “intelligence”. As bases de dados pesquisadas foram a Pubmed, Scielo, Scopus, Springer, IEEE, Web of Science (Figura 31).

Para a seleção dos artigos, foram considerados como critério de inclusão aqueles cuja leitura dos resumos referenciavam os algoritmos de IA aplicados ao xadrez, como também a tomada de decisão por estes algoritmos e pelos seres humanos.

Já os critérios de exclusão de artigos, após leitura dos resumos, foram daqueles que não se referiam ao JX ou que não abordavam a tomada de decisão de algoritmos de IA especialistas em xadrez e de seres humanos (Quadro 6). A Figura 32 configura o “status” dos artigo, a quantidade daqueles que foram aceitos, rejeitados, duplicados ou não classificados. E a Figura 33 apresenta a “nuvem de palavras” geradas pelos artigos.

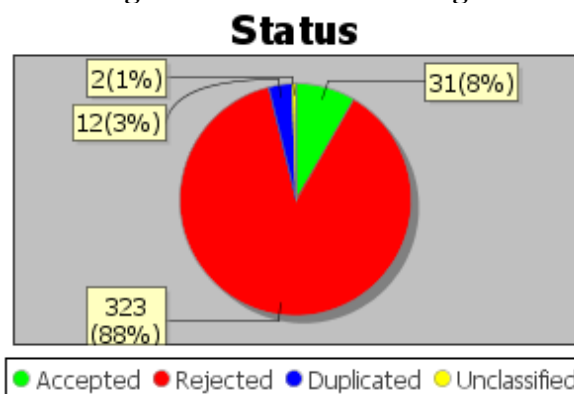
Figura 31 - Fontes de Pesquisa



Fonte: o autor, 2026

Total de artigos: 366 artigos (dois não foram importados)

Figura 32 - Status dos artigos



Fonte: o autor, 2026

Figura 33 - Nuvem de Palavras



Fonte: o autor, 2026

Quadro 6 - Artigos encontrados

Relevância	TÍTULO	AUTORES	COMENTÁRIOS	ANO
alta	Aplicacion de un modelo cognitivo de valoracion emotiva a la función de evaluacion de tableros de un programa que juega ajedrez	Laureano-Cruces, Ana Lilia <i>et al.</i>	Oferece oportunidade aos jogadores humanos opções de jogarem com uma "engine" mais fraca, em um nível mais próximo ao próprio jogador	2012
Sem relevância	In the Beginning of AI, There Were Games Playing Smart On Games, Intelligence, and Artificial Intelligence	Togelius, Julian <i>et al.</i>	Breve história de algoritmos, por exemplo, Minimax e árvore de Monte Carlo; primeiro capítulo do livro "Playing Smart"	2018
baixa	A Short History of Computer Chess	Marsland, T. A.	Sobre a história dos programas de xadrez.	1990
média	A general reinforcement learning algorithm that masters chess, shogi, and Go through self-play	Silver, David <i>et al.</i>	referência sobre o AlphaZero e sobre seu algoritmo	2018
alta	Aligning Superhuman AI with Human Behavior: Chess as a Model System	McIlroy-Young, R. <i>et al.</i>	modelagem e treinamento do pensamento humano acoplado ao AlphaZero para se prever os lances (e erros) humanos com a ideia de igualar o nível de jogo entre humano e engine. A pergunta é: "Que movimento um humano neste nível de habilidade jogará?".	2020
alta	Chess AI: Competing Paradigms for Machine Intelligence	Maharaj, Shiva <i>et al.</i>	Inspirando-se em como os humanos resolvem problemas de xadrez, perguntamos se as máquinas podem possuir uma forma de imaginação.	2022
alta	Chess knowledge predicts chess memory even after controlling for chess experience: Evidence for the role of high-level processes	Lane, David M. & Chang, Yu-Hsuan A.	Nas discussões são mostradas como os jogadores tem dificuldades em recuperar memórias que eles não entendem, as posições que não fazem sentido; comparação com engines	2017
Sem relevância	Chess Moves Prediction using Deep Learning Neural Networks	Panchal, Hitanshu <i>et al.</i>	Convolution Neural Network, e os treinos de máquina para jogar xadrez; aspectos históricos sobre a	2021

			evolução do xadrez; fala sobre os algoritmos Alpha-Beta e Minmax	
média	Chess-Playing Programs and the Problem of Complexity	Newell, A.; Shaw, J. C.; Simon, H. A.	grande valor histórico do início da IA e xadrez	1958
Sem relevância	Clever computers-why should <i>engineers</i> use AI?	Chung, P. & Inder, R.	sobre a importância da IA na engenharia	1990
Sem relevância	Data on human decision, feedback, and confidence during an artificial intelligence-assisted decision-making task	Chong, Leah <i>et al.</i>	trata-se de construção de um banco de dados que pode vir a ser consultado	2023
baixa	Direct Human-AI Comparison in the Animal-AI Environment	Voudouris, Konstantinos <i>et al.</i>	algumas referências sobre a comparação entre animais e a IA, tem algum valor	2022
alta	Expertise and Intuition: A Tale of Three Theories	Gobet, Fernand & Chassy, Philippe	boas referências sobre IH e decisões; intuição faz parte da expertise. Intuição e criatividade, analítica ou holística	2009
média	Expertise in complex decision making: The role of search in chess 70 years after de Groot	Connors, M.H. ; Burns, B.D. & Campitelli, G.	referências históricas sobre a tomada de decisões com humanos	2011
muito baixa	Flexible decisions and chess expertise	de Lafuente, V.	Sobre decisões flexíveis no xadrez	2011
alta	Human and computer preferences at chess	Regan, K.W. & Biswas, T. and Zhou, J.	comparação estilística de decisões; humanos se subestimam quando jogam com jogadores mais fortes, não gostam de fazer seus lances primeiro e uma cooperação entre homem e máquina é desejada	2014
média	Hybrid Optimization for Developing Human Like Chess Playing System	Chole, Vikrant & Gadicha, Vijay	comparações e argumentos sobre técnicas de algoritmos; escreve sobre as várias técnicas de algoritmos	2022
Sem relevância	Intrinsic chess <i>ratings</i>	Regan, K.W. & Haworth, G.M.	considerações sobre <i>rating</i> de jogadores	2011
alta	Is chess the drosophila of artificial intelligence? A social history of an algorithm	Ensmenger, Nathan	História social dos algoritmos de xadrez	2012
muito baixa	Judgment and Decision-Making: Experts and Novices Evaluation of Chess Positions	Eisele, Per	comparação entre humanos e suas preferências; diferenças entre mestres e amadores em suas escolhas. Comparação de <i>ratings</i> de humanos	2004
muito baixa	Monte Carlo tree search in Kriegspiel	Ciancarini, P. & Favini, G.P.	Abordagem sobre a método Monte Carlo para o jogo Kriegspiel.	2010
Sem relevância	On Computational Humor	Taylor, Julia M. & Mazlack, Lawrence J.	Abordagem do humor nos sistemas computacionais	2010

muito baixa	Observability of the mind the metodological inspiration provides by chess for digitalizing decision-macking process	Peng, Hongxia & Delorme, Axel	artigo curto sobre pesquisas de como funciona o reconhecimento de padrões pelo ser humano influenciou algoritmos de xadrez	2019
Sem relevância	Research on the Advantages and Equilibrium of Computer Game with Incomplete Information	Dong, Meiqi <i>et al.</i>	comparação em jogos de informação incompleta	2017
Sem relevância	Response time distributions in rapid chess: a large-scale decisionmaking experiment	Sigman, Mariano <i>et al.</i>	sobre a rapidez nas respostas de lances	2010
muito baixa	Search Versus Knowledge in Human Problem Solving: A Case Study in Chess	Bratko, Ivan, Hristova, Dayana & Guid, Matej	comparação entre conhecimento e busca de lances entre humanos	2016
média	Simulation-Based Algorithms for Markov Decision Processes: Monte CarloTree Search from AlphaGo to AlphaZero	Fu, Michael C.	como funcionam os algoritmos Alpha-Beta (Minimax) e Monte Carlo Tree Search	2019
alta	Six Suggestions for Research on Games in Cognitive Science	Chabris, Christopher F.	várias ideias importantes de pesquisa e comparações entre decisões humanas e de máquina	2017
média	Slow and fast chess performance across three expert levels	Blanch, A.; Ayats, A. & Cornadó, M.P.	comparação das mudanças entre o pensamento lento e rápido (Khanemann) em tres níveis de expertise dos jogadores	2020
baixa	The role of intuition and deliberative thinking in expert's superior tactical decision-making	Moxley, Jerad H. <i>et al.</i>	o papel do treinamento em xadrez em detrimento do papel da "intuição" (heurística);(Khanemann)	2012
Sem relevância	The tell-tale heart: heart rate fluctuations index objective and subjective events during a game of chess	Leone, Maria J. <i>et al.</i>	sobre os batimentos cardíacos (emoções) durante as tomadas de decisões	2012

Fonte: o autor, 2026

O Quadro 7 apresenta a relevância dos artigos encontrados:

Quadro 7 - Relevância dos artigos

RELEVÂNCIA	QUANTIDADE
SEM RELEVÂNCIA	9
MUITO BAIXA	5
BAIXA	3
MÉDIA	5
ALTA	8
TOTAL	31

Fonte: o autor, 2026

O resultado da revisão sistemática nas bases de dados trouxe 31 artigos, porém apenas 5 obedeceram aos critérios de inclusão e exclusão e foram considerados de

relevância para este estudo, sobre as diferenças nos processos de tomada de decisão entre os algoritmos de IA, IH e o JX. O resultado da pesquisa se encontra no Quadro 8.

Quadro 8 - Artigos incluídos

TÍTULO	AUTORES	ANO
Aplicacion de un modelo cognitivo de valoracion emotiva a la función de evaluacion de tableros de un programa que juega ajedrez	Laureano-Cruces, Ana Lilia, Hernández-González, Diego Enrique, Mora-Torres, Martha, Ramírez-Rodríguez, Javier	2012
Aligning Superhuman AI with Human Behavior: Chess as a Model System	McIlroy-Young, R. and Sen, S. and Kleinberg, J. and Anderson, A.	2020
Chess AI: Competing Paradigms for Machine Intelligence	Maharaj, Shiva and Polson, Nick and Turk, Alex	2022
Chess knowledge predicts chess memory even after controlling for chess experience: Evidence for the role of high-level processes	Lane, David M. and Chang, Yu-Hsuan A.	2018
Human and computer preferences at chess	Regan, K.W. and Biswas, T. and Zhou, J.	2014

Fonte: o autor, 2026

13.2 A Humanização do Algoritmo de Xadrez

Os lances executados pelo algoritmo AlphaZero, muitas vezes incompreensíveis por adotarem decisões incomuns, dificultam as possibilidades de se aprender as novas ideias advindas de seus movimentos no JX. Na tentativa de “humanizar” o AlphaZero, McIlroy-Young *et al.* (2020), aplicaram ao algoritmo tomadas de decisões mais parecidas com as dos humanos. A técnica usada foi a de aprender padrões com milhões de jogos entre humanos jogados *online* e assim tornar as decisões do computador em dado contexto de jogo como um padrão mais humano e assim conseguir prever com mais acerto as decisões que os humanos tomariam durante as partidas. A tentativa parece estar em acordo de regular o AlphaZero para que possa colaborar com o ser humano nas tomadas de decisão. Ao final, os algoritmos conseguiriam ajudar o ser humano em seus aprendizados para jogar xadrez (McIlroy-Young *et al.*, 2020). Esta é uma tentativa de aproximar a força da *engine* (algoritmos especializado em jogar xadrez) com o nível de força do jogador humano que está jogando. Em outras palavras, buscar o viés de comportamento pela média de força do humano (em uma dada faixa de força do jogador, a sua pontuação no *ranking* dos enxadristas, ou *rating*). Porém, este aprendizado não leva em conta as

características individuais dos jogadores, suas raízes, culturas. Isto acarreta apenas um padrão da média, não a característica individual de cada jogador de xadrez.

A partir deste artigo parte a ideia das questões culturais dos diagramas 1 e 2 do *survey* (capítulo 14).

13.3 Atribuição de Emoções aos Algoritmos de Xadrez

Atribuir emoções às *engines*, na tentativa de simular os sentimentos dos jogadores humanos nas tomadas de decisão de lances pela *engine*, foi pesquisada por Laureano-Cruces *et al.* (2012). Os autores deram nome à *engine* com a alcunha “*Deep Feeling*” e argumentaram sobre a importância do estado emocional dos jogadores humanos durante o processo de tomadas de decisão nas partidas de xadrez em uma linha de investigação chamada de “computação afetiva”. Os resultados obtidos foram promissores (Laureano-Cruces *et al.*, 2012).

Este artigo inspirou os diagramas 4 e 5 sobre ToM e “fair play” do *survey* (capítulo 14).

13.4 Análises de Algoritmos e de Humanos em Posições de Xadrez

Uma das possibilidades de comparar as decisões humanas com as IAs são com os finais de partida de xadrez, quando poucas peças estão no tabuleiro. Uma posição de final de jogo famosa pela extrema dificuldade em se resolver o problema pelos jogadores humanos, composta por um enxadrista amador alemão no fim dos anos 1970, mas que se popularizou com o grande mestre inglês Jim Plaskett em um torneio de xadrez realizado em 1987, só foi resolvida, na ocasião, pelo ex-campeão mundial Mikhail Tal (Maharaj, Polson e Turk, 2022), conhecido pela sua imaginação e fantasia férteis, considerado um dos maiores neste quesito (Kasparov, 2004). A composição foi aplicada a mestres de xadrez durante o torneio de 1987. Os autores do estudo compararam duas *engines* de xadrez, o “*Stockfish 14*”, com seu algoritmo “Minimax” e o algoritmo “*Leela Zero*”, baseado em RNA. Os autores descobriram que o algoritmo Minimax se saiu melhor que a rede neural artificial, ou seja, o Minimax foi mais eficiente para o xadrez do que a rede neural artificial, naquele dado momento (Maharaj, Polson e Turk, 2022).

Este artigo colaborou com a ideia do diagrama 7 do *survey* sobre comparar humanos e *engines* de xadrez em posições complexas (capítulo 14).

13.5 Conceitos sobre o Xadrez e o Reconhecimento de Padrões

Uma controvérsia sobre o reconhecimento de padrões pelos mestres de xadrez está ligada ao estudo intenso do jogo (De Groot, 1988; Chase e Simon, 1973; Gobet, 2019). Esta habilidade está correlacionada com a memória e com a aprendizagem dos conceitos enxadrísticos, ou seja, o alto nível de conceitualização aliado a um alto reconhecimento de padrões pelos jogadores de xadrez atuam de forma importante na memorização do jogo. Mais exatamente, quanto maior os conceitos a respeito do jogo, é mais fácil o aprendizado do reconhecimento de padrões, por sua consequente melhor memorização. Isso implica que, para tomar decisões, a conceitualização é importante para o posterior reconhecimento de padrões pelos seres humanos (Lane e Chang, 2018). Isto poderia explicar o sucesso do poder computacional das *engines* e a não necessidade de reconhecer padrões para as suas tomadas de decisão e mesmo para os seres humanos, tornando o conhecimento conceitual sobre o jogo mais importante do que simplesmente ter na memória inúmeros padrões e “*chunks*”.

Nas discussões do artigo de Lane e Chang (2018), é mostrado como os jogadores têm dificuldades em recuperar memórias que eles não “entendem”, não tem familiaridade com o conceito, as posições que não fazem sentido e, como consequência, os jogadores não reconhecem os padrões. Os autores concluem que é mais importante entender uma posição do que reconhecer padrões. Nas teorias e modelos de memória para se jogar xadrez deveriam ser incluídas as diferenças individuais das habilidades cognitivas, para se averiguar o quanto a habilidade cognitiva influi na memorização do xadrez para os seres humanos (Lane e Chang, 2018).

Este artigo colaborou com as ideias de experimento sobre reconhecimento de padrões com humanos, diagrama 6 do *survey*, capítulo 14.

13.6 Avaliações Humanas são Relativas e de *Engines* Concretas

Regan, Biswas e Zhou (2014) fizeram análises estatísticas em bases de dados de xadrez e descreveram gráficos que apoiam a hipótese de que os humanos percebem as diferenças das posições de xadrez de acordo com avaliações relativas, enquanto as *engines* só avaliam de forma concreta, através dos cálculos de seus algoritmos. Ou seja, os humanos são afetados por vários fatores, como por exemplo a questão subjetiva de

enfrentar um adversário mais forte, fazendo avaliações piores do que fariam com jogadores mais fracos.

Os autores destacam também que há uma desvantagem significativa para os enxadristas em terem que escolher primeiro um lance do que se fosse a vez de seu adversário, o que se julgou um pessimismo em suas avaliações, enquanto as *engines* preferem “esperar para ver” (Regan, Biswas e Zhou, 2014.) É possível, no entanto, que a cooperação humano-computador produza resultados melhores do que qualquer ação separada de máquina e ser humano.

13.7 Discussões sobre a Revisão de Artigos

A partir do JX é possível extrair exemplos a respeito das diferentes perspectivas entre a IA e a IH e a capacidade que os humanos têm de entender, abstrair, se emocionar. Percebe-se uma tendência em “humanizar” os algoritmos, ou seja, inserir emoções na IA.

Os relatos sobre campeonatos mundiais de xadrez estão repletos de situações em que os jogadores tomam decisões fortemente impactadas pelas suas emoções. Assim aconteceu no campeonato mundial de 1950 entre os russos Botvinnik e Bronstein (Kasparov, 2004). Em vários momentos do confronto de melhor de vinte e quatro partidas entre ambos, lances que, normalmente seriam fáceis de calcular por jogadores de tão alto nível, foram afetados pelas emoções e decisões frágeis acabaram por prejudicar ambos os jogadores. Botvinnik, preferido pelo poder político da então União Soviética, desconfiava inclusive de seus analistas particulares. Bronstein, judeu, tinha seu filho vigiado pela polícia política (Kasparov, 2004). Um algoritmo para jogar xadrez do nível destes mestres seria afetado por estas emoções? Cometeriam as mesmas falhas? Seria possível implementar nestes algoritmos estas emoções e com qual propósito?

Outro exemplo, talvez o mais emblemático das disputas pelo campeonato mundial de xadrez, aconteceu entre o estadunidense Bobby Fischer e o então campeão do mundo, o soviético Boris Spassky. Este confronto ficou muito famoso ao redor do planeta, pois tratava-se de um embate entre Estados Unidos e União Soviética, em plena guerra fria, pela supremacia em um jogo, o xadrez, cujo paradigma é o da inteligência (Kasparov, 2006).

Fischer, na primeira das vinte e quatro partidas que seriam disputadas pelo campeonato mundial, fez um lance que surpreendeu a todos, um erro estratégico de

amadores. E assim, Spassky venceu a primeira partida (ver capítulo 14, seção 14.6). Na segunda partida, o mais inusitado aconteceu. Fischer simplesmente se recusou a disputar a partida e não compareceu ao salão de jogos (Kasparov, 2006).

Estas duas atitudes, entregar as duas primeiras partidas para o adversário, muito estranhas para um jogador que possuía uma enorme ambição pelo título mundial, é ainda hoje tema controverso. Muitos analistas concordam em dizer que foi uma atitude deliberada do americano para abalar psicologicamente o adversário (Kasparov, 2006). Ou seja, poderia ter sido uma estratégia sem precedentes na história. A pergunta que fica é, poderia uma IA simular tal situação? A IA entenderia que este tipo de comportamento poderia ser uma estratégia, pois foge aos padrões de lógica, cálculos e probabilidades presentes em seus algoritmos? Ou a IA simplesmente não precisa deste tipo de estratégias para chegar ao objetivo, qual seja, vencer o oponente?

Com as leituras foi possível conhecer as técnicas de tomada de decisão da IA, ou seja, o Minimax Alpha-Beta (no caso do *Deep Blue* da IBM e os softwares atuais que já superaram o ser humano no JX, como é o caso do Stockfish) e o DRL (no caso do AlphaZero da Deep Mind).

O ser humano, ao contrário das *engines*, calcula apenas algumas possibilidades de lances, de acordo com seus conhecimentos técnicos que foram aprendidos para se jogar xadrez. Por isso o ser humano “entende” a posição e o que deve ser selecionado como “lances candidatos”, que implica em não vasculhar toda a árvore de decisão e sem a necessidade de aprofundar cálculos (obviamente, dependendo da posição no tabuleiro). O ser humano “entende” a posição, o que as *engines* não conseguem fazer, pois são apenas cálculos (ver capítulo 14, seção 14.8 e 14.9).

Ribeiro-Doggers (2025), ao perguntar a Matthew Sadler (Sadler e Regan, 2019) pesquisador sobre *engines* de xadrez com redes neurais, se as *engines* modernas “entendiam” o jogo de xadrez, Sadler foi categórico em dizer que não entendem nada e muito menos sabem que estão jogando xadrez, mas parecem que sabem. Isso se deve ao estilo de jogo de AlphaZero e LC0, pois as redes neurais estimam lances que maximiza a mobilidade das peças, a atividade, a dinâmica. Este ponto é uma revolução na forma de pensar o xadrez, é muito mais importante que as peças tenham muitas opções de jogadas a serem feitas, o que amplia a sua força relativa.

14. RESULTADOS DA PESQUISA SURVEY

14.1 Perfil dos Participantes

O *survey* obteve amostra de 84 participantes que jogam xadrez em vários níveis de força (*rating*), com o objetivo de avaliar características sobre funções cognitivas (Quadro 9). Dos 84, apenas 1 desistiu de participar da pesquisa.

Os diagramas apresentados aos participantes do *survey*, de 1 a 8, foram elaborados para avaliar características da tomada de decisão dos jogadores de xadrez, dentre as quais, a flexibilidade mental, atenção seletiva, ToM, reconhecimento de padrões, cultura e estética, abstração.

Quadro 9 – Diagramas e funções cognitivas

Diagrama	Avaliar
1	Flexibilidade, cultura, estética;
2	Flexibilidade, cultura, estética;
3	Atenção seletiva; percepção visual;
4	ToM, entender o contexto;
5	ToM, entender o contexto;
6	Reconhecimento de padrões, memória semântica;
7	Memória de trabalho, raciocínio lógico, entender a posição;
8	Abstração; entender a posição;

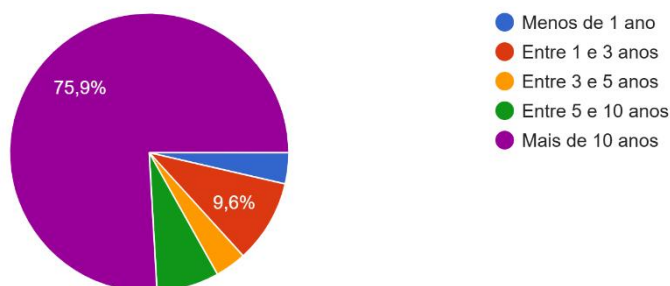
Fonte: o autor, 2026

O Gráfico 1 mostra que as pessoas que responderam ao questionário de oitenta e três participantes, 75,9% (63 participantes) têm mais de dez anos de experiência com o xadrez; 7,2% (6 participantes) têm entre 5 e 10 anos de experiência; 3,6% (3 participantes) têm entre 3 e 5 anos de experiência; 9,6% (8 participantes) têm entre 1 e 3 anos de experiência e 3,6% (3 participantes) têm menos de um ano de experiência em jogar xadrez.

Gráfico 1 - Tempo que pratica o xadrez

Há quanto tempo você joga xadrez?

83 respostas



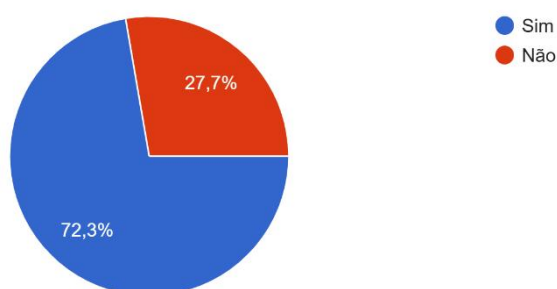
Fonte: o autor, 2026

O Gráfico 2 mostra quais participantes são mais ativos e competitivos ao participarem de torneios (eventos). Os que responderam que já jogaram alguma vez foi de 72,3% com 60 participantes e 27,7% com 23 participantes nunca jogaram. Com este resultado, é possível saber que ao menos 60 participantes têm ótimas noções de xadrez.

Gráfico 2 - Participação em Torneios de Xadrez

Você já jogou ou joga torneios de xadrez?

83 respostas



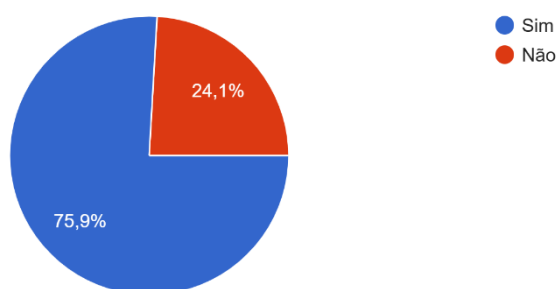
Fonte: o autor, 2026

No Gráfico 3, 63 participantes, 75,9 %, já leram livros de xadrez, confirmando a amostra de jogadores experientes com mais de 10 anos de prática com o xadrez.

Gráfico 3 - Leitura de Livros

Já leu algum livro de xadrez?

83 respostas



Fonte: o autor, 2026

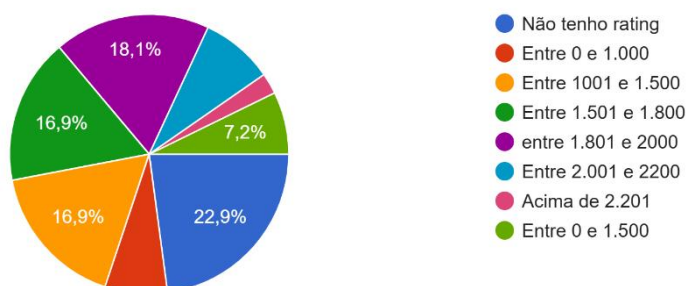
O Gráfico 4 apresentou uma distribuição de força de jogo (*rating*) bastante heterogênea, abrangendo ampla diversidade. Apenas 2 participantes, 2,4% da amostra, possuem *rating* acima de 2.200 pontos, que corresponde a uma categoria de mestre; 7 participantes, 8,4% tem *rating* entre 2.001 e 2.200 pontos, que corresponde a uma

categoria de mestre nacional; 15 participantes, 18,1% tem entre 1.801 e 2.000 pontos de *rating*, o que corresponde a um jogador de 1ª categoria ou candidato a mestre nacional; 14 participantes, 16,9% estão entre 1.501 e 1.800 pontos de *rating*, o que corresponde a um jogador de 2ª categoria, jogador de clube; 14 participantes, 16,9% se disseram com *rating* entre 1.001 e 1.500 pontos de *rating* o que corresponde a um jogador de 3ª categoria, frequentador de clube; 12 participantes, 14,4% tem entre 0 e 1.500 pontos, considerado um jogador pouco ativo; finalmente, 19 participantes, 22,9% disseram não possuir *rating*, caracterizando-os como jogadores iniciantes. Dos 83 participantes, 45 (54,2%) são jogadores abaixo de 1500 pontos, caracterizando a amostra com maioria de amadores pouco atuantes.

Aqui é necessário deixar claro que não houve neste trabalho de tese uma equalização dos diferentes tipos de *rating* dos enxadristas que responderam à pesquisa (*rating* FIDE, Chess.com, Lichess).

Gráfico 4 - Rating

Qual o seu rating (FIDE, CBX, Lichess, Chees.com, etc.)
83 respostas



Fonte: o autor, 2026

14.2 Diagrama 1, a Escolha Estética do Xeque-mate

Na Figura 34 (Diagrama 1 do *survey*), todas as opções de lances oferecidas aos participantes conduzem à posição de xeque-mate. A diferença entre as opções se refere ao gosto pessoal de cada jogador.

O Gráfico 5 apresenta o percentual de escolha de cada lance pelos participantes. Por exemplo, o “xeque-mate a descoberto”, o movimento do Rei branco que permite a Torre de h1 realizar o xeque-mate, é raríssimo na prática do xadrez, tem um viés estético

de alto valor. O lance Rei branco para “g1” recebeu 3,6% das escolhas, enquanto Rei branco para “g3” recebeu 15,7%.

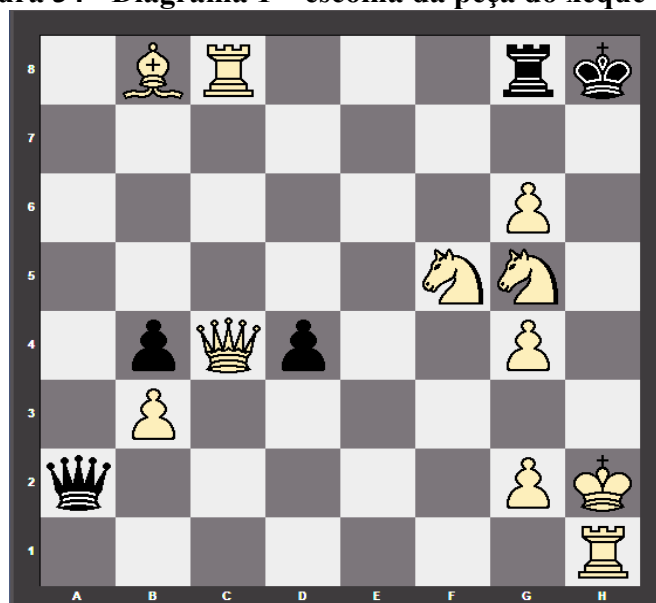
Já o lance Dama captura Torre preta em g8 (43,4% das respostas) é o lance típico do computador, da “força bruta”, onde a maior pontuação avaliada pela *engine* é a segunda a ser escolhida. Nota-se uma preferência dos jogadores sem *rating* ou com *rating* abaixo de 1.500 (24% dos 43,4%), ou seja, mais da metade daqueles que optaram pelo lance do computador, o que pode indicar pouco conhecimento estético/cultural sobre xadrez. O lance do Stockfish 18, Torre branca captura a Torre preta em “g8” teve 8,4% das preferências, outro lance da “força bruta”.

Outros participantes preferem o xeque-mate com o movimento do Cavalo (18,1%); com o movimento do Peão (7,2%), a peça mais humilde do jogo (simbolicamente, representou na Idade Média, camponeses e artesãos).

O xeque-mate de Bispo em “e5”, que esteticamente possa ser o menos preferível, teve 3,6% das preferências, enquanto o lance da Dama capturando Peão preto de “d4”, terceiro lance preferido pelo Stockfish, do ganho de material e da força bruta, também pouco estético, não teve sequer um voto.

Muitos significados podem ser extraídos do diagrama da Figura 34: a diversidade da escolha, as preferências estéticas por lances incomuns; porém houve a prevalência da escolha dos participantes pelos lances de “menor conhecimento” de xadrez, a decisão puramente materialista do cálculo frio do Stockfish 18 (Figura 35).

Figura 34 - Diagrama 1 – escolha da peça do xeque-mate

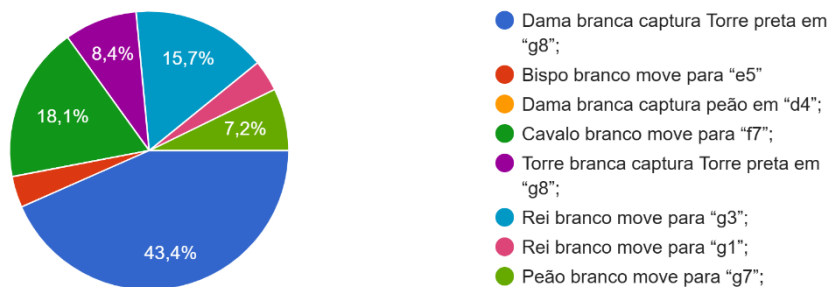


Fonte: o autor, 2026

Gráfico 5 – Percentual da Escolha do Lance pelo Participante

Diagrama 1 - Qual o lance de sua escolha?

83 respostas



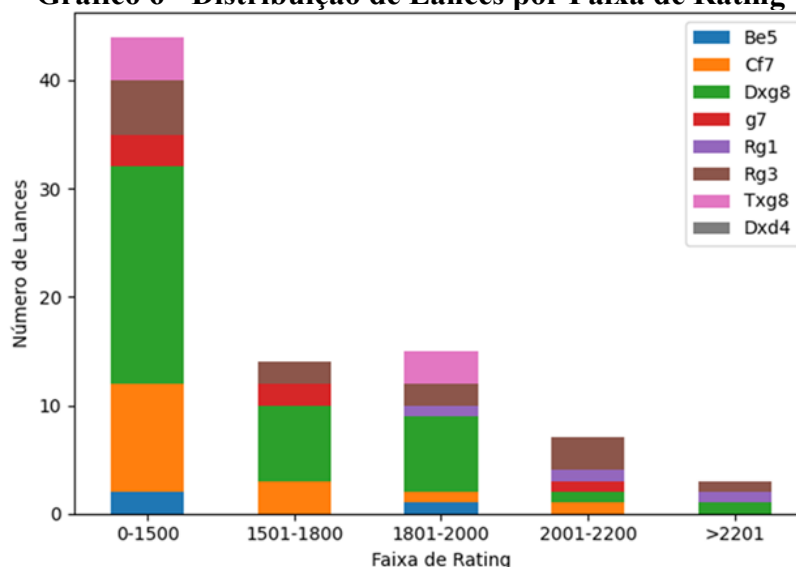
Fonte: o autor, 2026

Figura 35 - As cinco preferências do Stockfish no diagrama 1



Fonte: sítio www.lichess.com, construído pelo autor, 2026

No Gráfico 6 ficam claras as preferências de cada faixa de *rating* correlacionada com a força dos jogadores.

Gráfico 6 - Distribuição de Lances por Faixa de Rating

Fonte: o autor, 2026

Nas faixas de *rating* mais baixas (0-1500 e 1501-1800), o lance Dxg8 é a escolha preferencial da maioria. Isso indica um lance da “força bruta”, da capitalização de material.

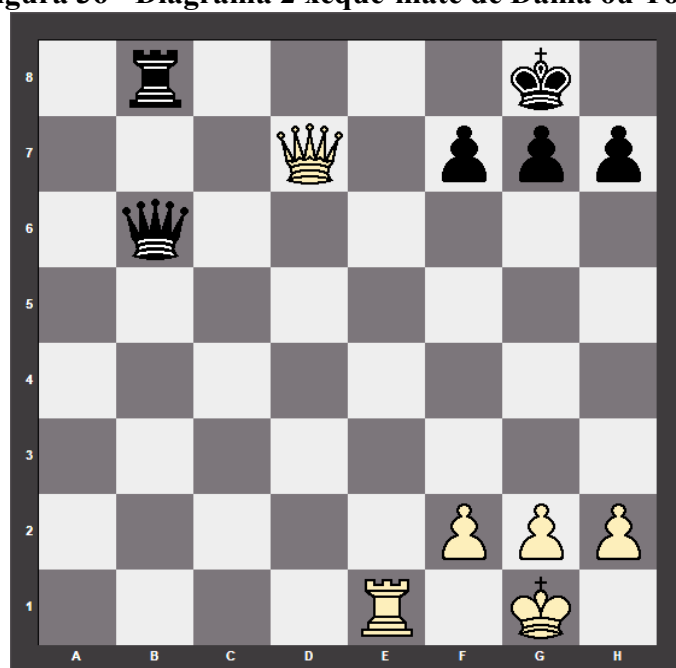
À medida que o *rating* aumenta (especialmente a partir de 2001), a preferência por Dxg8 diminui drasticamente, e lances como Rg1 e Rg3 (os lances incomuns, muito raramente vistos nos xeque-mates) ganham espaço. Isso sugere que jogadores mais fortes identificam nuances na posição que os jogadores de *rating* menor ignoram.

O lances Be5 talvez seja o mais atípico, porém tem seu valor estético, pois é o caminho para trás na diagonal; já o mate de Peão (g7) é o lance de xeque-mate da peça mais humilde do jogo; ambos tiveram pouquíssima adesão em todos os níveis, sendo quase ignorados pela amostra. O terceiro lance da engine, Dxd4, sequer foi mencionando pelos participantes.

14.3 Diagrama 2, o Sacrifício de Dama

Na Figura 36 (Diagrama 2 do *survey*), ocorre a mesma situação da figura 33. Os jogadores mais experientes (41%) preferem o lance “mais bonito” esteticamente, sacrificar a Dama para dar xeque-mate no lance seguinte, ao contrário de jogadores com pouco *rating* e do Stockfish 18 (Figura 37), que preferem permanecer com o material, lance Torre para “e8” e não sacrificar a Dama em “e8” primeiro. O que buscaram os participantes, o lance prosaico da segurança e a certeza, ou o lance artístico e despojado?

Figura 36 - Diagrama 2 xeque-mate de Dama ou Torre

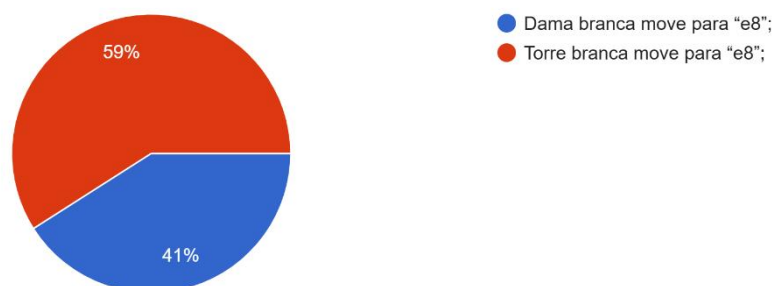


Fonte: o autor, 2026

Gráfico 7 - Percentual de Escolha do Participante

Diagrama 2 - Qual o lance de sua escolha?

83 respostas



Fonte: o autor, 2026

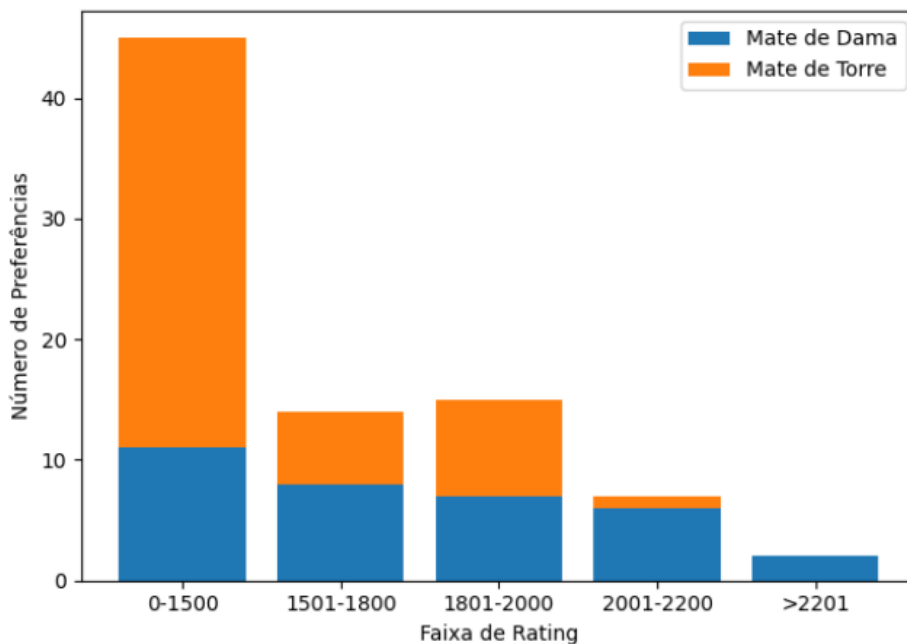
Figura 37 - Decisão do Stockfish 18 – Primeiro a Torre



Fonte: sítio www.lichess.com, construído pelo autor, 2026

No Gráfico 8 é verificada a associação entre a força dos jogadores (*rating*) e a preferência pelo sacrifício da Dama ou da Torre.

Gráfico 8 - Sacrifício de Dama ou Torre?



Fonte: o autor, 2026

Nota-se que, na faixa de *rating* entre 0 e 1500 há uma forte preferência pelo mate de Torre (34 contra 11). Na faixa de *rating* de 1500-1800, a distribuição se equilibra. Na

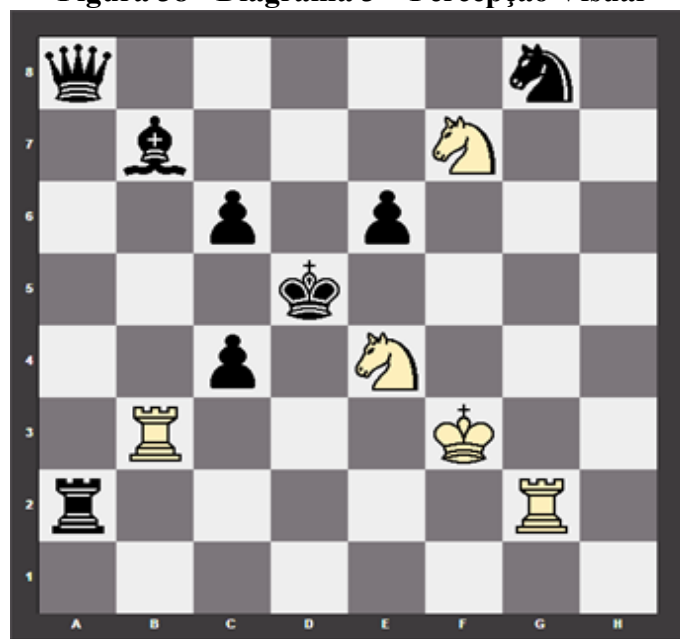
faixa de *rating* de 1801-2000, permanece o equilíbrio. Nas duas maiores faixas de *rating*, 2001-2200 e maior que 2201, há uma clara preferência pelo xeque-mate com o sacrifício de Dama. Ou seja, jogadores menos experientes parecem preferir o xeque-mate sem o risco de perder material e jogadores mais fortes tem a preferência do padrão estético/cultural, qual seja, o sacrifício da Dama para dar xeque-mate no lance seguinte. Em outras palavras, observa-se uma associação muito forte entre a expertise e a percepção estética/cultural.

14.4 Diagrama 3, Percepção Visual

Na Figura 38 (Diagrama 3 do *survey*), a proposta é procurar saber o que pensam os jogadores ao ver o diagrama 3. Querer resolver o problema de xadrez ou simplesmente “enxergar” (imaginar) a letra “X”. A pergunta proposta foi “qual a primeira ideia que lhe vem à cabeça?”. A ideia deste diagrama foi de observar a relação de percepção visual do tabuleiro em si ou se os jogadores têm atenção seletiva (“focada”) exclusivamente nas jogadas de xadrez.

Nesta questão, apenas 13% dos participantes viram o “X”, ou uma simetria, “beleza”. Os demais tentaram jogadas como Ta2 (9,6%), lances de Cavalo (15,6%), lance de Torre para “g5” com xeque ao Rei preto (22,9%). Outros participantes mencionaram a segurança dos Reis, a possibilidade de empate ou a busca por um xeque-mate. Na Figura 39, temos quais lances são considerados pelo Stockfish 18.

Figura 38 - Diagrama 3 – Percepção Visual



Fonte: o autor, 2026

Figura 39 - Decisão do Stockfish 18 – Cc3+ com empate



Fonte: sítio www.lichess.com, construído pelo autor, 2026

No Gráfico 9 é observada a associação entre a força do jogador e a percepção visual, a atenção seletiva.

Gráfico 9 - Atenção Seletiva, percepção visual

Fonte: o autor, 2026

Ao contrário das perguntas sobre cultura e estética, a percepção visual do "X" foi muito baixa em todas as faixas de *rating*. Curiosamente, a faixa 0-1500 foi a que teve o maior número absoluto de participantes que notaram o "X" (4 pessoas), possivelmente por estarem olhando o tabuleiro de uma forma menos seletiva ("focada"), ou atentas apenas na tática e mais no jogo em geral. Nas faixas acima de 2001, ninguém (0%) relatou ter visto a letra "X". Isso pode indicar que jogadores mais avançados ficam tão seletivos nas relações de força entre as peças, que padrões puramente estéticos ou geométricos irrelevantes para o jogo passam despercebidos.

Nas faixas de *rating* mais alto (2001-2200 e >2201), ninguém (0%) percebeu a letra "X". Isso sugere que jogadores de maior força tem a sua atenção voltada para a função tática das peças e ignoram padrões geométricos puramente visuais.

A maior taxa de percepção (proporcionalmente) ocorreu na faixa intermediária (1501-2000), onde cerca de 20% dos jogadores notaram o "X" "desenhado" no tabuleiro.

Independentemente da força de *rating*, a maioria dos enxadristas não nota padrões que não influenciam diretamente no jogo, possuem atenção seletiva para as possibilidades de lances. A maioria das respostas tentou achar alguma jogada ganhadora ou de xeque-mate.

14.5 Diagrama 4, “*Fair Play*”, o Contexto

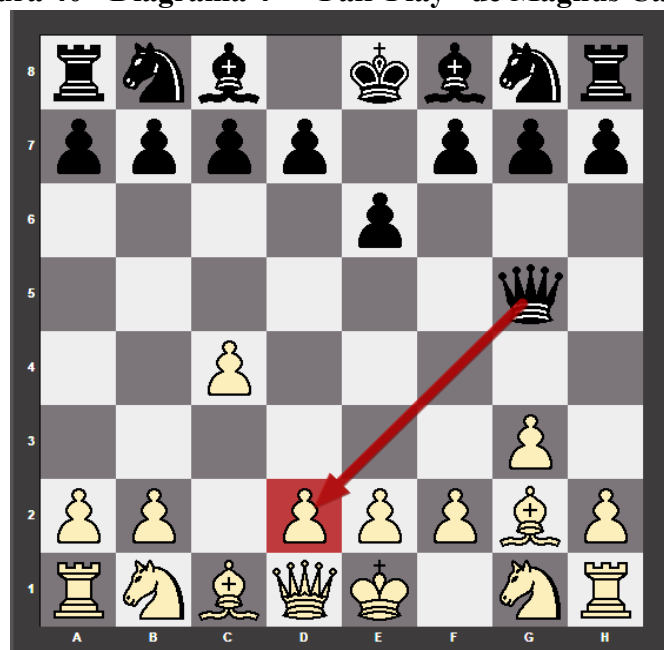
Na Figura 40 (Diagrama 4 do *survey*), foi apresentado o seguinte texto: “O Diagrama 4 apresenta um lance (seta vermelha) feito por Magnus Carlsen, ex-campeão mundial (com as peças pretas) contra Ding Liren, atual campeão mundial (que jogava de brancas), em uma partida *online*. Você consegue explicar, em poucas palavras, o porquê de Carlsen ter realizado este lance?”. Jogadores experimentados sabem que tal situação não condiz com a força e o conhecimento de um jogador do nível de campeão mundial.

Em 2021, na Copa do Mundo de Xadrez online, Magnus Carlsen, então campeão do mundo, recebeu um prêmio pela sua conduta de “*fair play*”. Isto porque, na primeira partida, que estava em posição de empate, o provedor de internet do computador de Ding desligou, “caiu”, não sendo possível que a partida fosse concluída e por isso, Carlsen foi declarado vencedor. Seria possível que uma *engine* de xadrez pudesse entender e realizar um “*fair play*”? Logo no terceiro lance, Carlsen jogou um “sacrifício” de Dama, inexplicável sem o contexto descrito acima. Uma tomada de decisão em que o norueguês preferiu devolver o ponto ganho da partida anterior, que não teve o mérito devido. Uma atitude cavalheiresca, talvez inusitada. Uma decisão individual, em que vários fatores, culturais, sociais, esportivos, dentre outros puramente humanos, determinou o empate no confronto, pois Carlsen abandonou a partida em seguida ao lance de Ding capturar a Dama que lhe fora entregue. Os espectadores da partida e comentaristas “entenderam” o que havia ocorrido com o lance extravagante de Carlsen. Dentre as possibilidades de lance para as pretas, tal lance é o último da árvore de decisão da *engine*.

A ideia de pesquisa com este diagrama foi verificar, entre os participantes, sobre tentar se colocar no lugar de Carlsen, ou seja, descobrir o que uma pessoa está pensando, empatizar, a função cognitiva da ToM, porém, sem fornecer o contexto da história.

Entender o contexto, mesmo sem ter sido mencionada toda a história, pode revelar uma das diferenças entre humanos e máquinas. Neste caso, apenas 19% dos participantes ou sabiam o que tinha acontecido ou perceberam que foi uma atitude de Magnus Carlsen para perder a partida. Os demais simplesmente não conseguiram entender a situação ou consideraram situações que não condiziam com o jogo de alto nível, que foram 81% dos participantes.

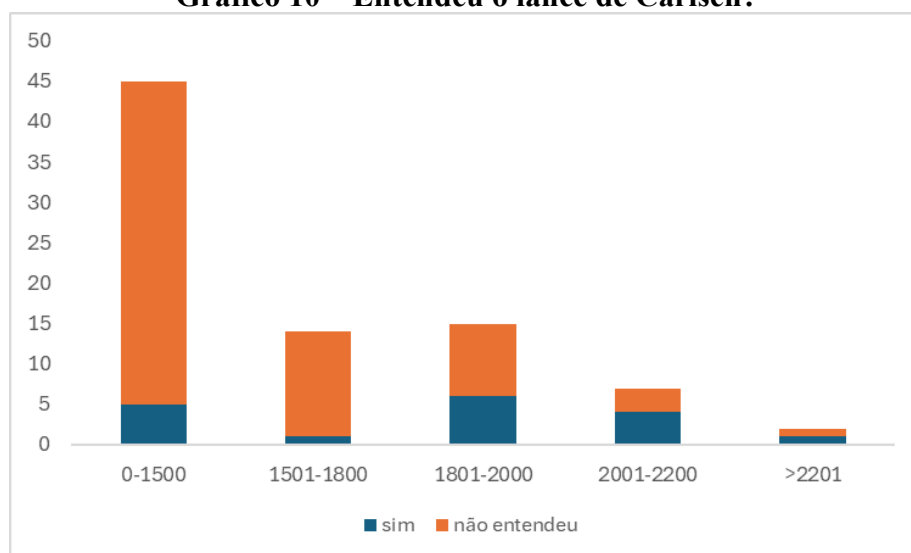
Figura 40 - Diagrama 4 – “Fair Play” de Magnus Carlsen



Fonte: o autor, 2026

O Gráfico 10 apresenta a associação de força dos participantes com o entendimento do porquê Magnus “entregou” a Dama e a partida para Ding Liren.

Gráfico 10 – Entendeu o lance de Carlsen?



Fonte: o autor, 2026

Na faixa de *rating* entre 0-1500, a grande maioria (89%) "não entendeu" a situação do “*fair play*”, o que é natural para jogadores iniciantes que ainda não acompanham o cenário profissional de perto. A partir dos 1801 pontos de *rating*, a percepção muda drasticamente. O número de participantes que conhecia ou reconheceu a atitude de Carlsen em “entregar” a partida cresce significativamente. Nas faixas de elite (acima de

2001), o entendimento sobre as atitudes de Magnus Carlsen é muito mais equilibrado e presente, refletindo um maior acompanhamento das notícias do mundo do xadrez.

14.6 Diagrama 5, Bobby Fischer Errou?

A Figura 41 (Diagrama 5 do *survey*), apresenta um dos momentos mais controversos da história do xadrez. A posição é da primeira partida do campeonato mundial de 1972 entre o então campeão mundial, o russo/soviético Boris Spassky e o desafiante ao título, o estadunidense Bobby Fischer. A Figura 42 mostra as análises do Stockfish 18, com clara posição de empate.

O encontro pelo campeonato mundial, em plena “Guerra Fria”, ficou conhecido como o “*match* do século”. Simbolicamente, era o antagonismo entre o mundo comunista e capitalista, uma disputa pela superioridade da “inteligência” dentro de um jogo tão simbólico e representativo, o xadrez. Políticos como Henry Kissinger (então conselheiro de segurança nacional dos EUA) conversaram pessoalmente com Fischer motivando-o para o campeonato mundial ocorrido em Reykjavík, na Islândia.

A ideia de pesquisa deste diagrama foi o de verificar entre os participantes se eles se colocariam no lugar de Fischer (aqui, mais uma vez, a ToM), porém, desta vez, com o relato do contexto histórico de partida tão importante. Outro ponto de investigação da pergunta foi o de conhecer a percepção dos participantes sobre a IA, se ela (as *engines* de xadrez, neste caso) são reconhecidamente mais fortes que os seres humanos no JX. Por fim, foi analisado se os participantes percebem estratégias em jogar lances que não são bons, apenas para desestabilizar o adversário, provocando vários tipos de reação, tanto do oponente, quanto do público, trazendo uma instabilidade emocional ao confronto. Vale lembrar que a partida seguinte a esta derrota “deliberada” de Fischer, também foi perdida pelo estadunidense, que sequer compareceu para jogar e perdeu por abandono. Mesmo assim, com estas duas derrotas iniciais, Fischer venceu com uma boa margem de vitórias contra Spassky, 12,5 pontos à 8,5 (Kasparov, 2006).

A pergunta feita foi a seguinte: “O Diagrama 5 é da primeira partida pelo campeonato mundial de 1972, entre o russo Boris Spassky e o desafiante, o estadunidense Bobby Fischer. Fischer, com as peças pretas, capturou o Peão branco em “h2” com seu Bispo de “d6”. Os analistas da época comentaram sobre “o incrível erro” cometido por Fischer, que veio a perder a partida, um lance controverso até os nossos dias. Você acha que uma IA poderia cometer tal erro? O lance de Fischer foi proposital? O que você pensa

a respeito?” A pergunta, aceitando respostas abertas, indicou, em primeiro lugar, que a IA “jamais cometeria esse erro”; houve também aqueles que admitem a possibilidade de erros das *engines* e outras respostas divergentes não conseguiram opinar. Ou seja, existe uma confiança para as análises e cálculos das *engines* de xadrez muito alta nos dias de hoje. Em outras palavras, as decisões da *engine* são altamente confiáveis, se tornaram o paradigma da força dos computadores.

As atuais *engines* de xadrez, como o Stockfish 18 (Figura 43), avaliam a posição com uma vantagem muito ligeira para as peças brancas (de Spassky, na referida partida), ou seja, o lance que Fischer jogou não tinha uma vitória clara para Spassky, foi um lance muito arriscado, mas que ainda podia sinalizar uma posição de empate no jogo. O “erro” em si (ou a intenção deliberada do “erro”), foi mínimo por parte de Fischer, ainda que um erro.

Por outro lado, uma parcela dos analistas do confronto à época entendeu que o lance de Fischer foi pensado, proposital, e não um erro. Foi uma escolha psicológica, para demonstrar a Spassky que não haveria empates fáceis e sim muita luta, mesmo que para isso o risco fosse muito elevado. Uma decisão humana, compreensível, de jogo. Uma perspectiva inviável para as *engines* de xadrez.

Figura 41 - Diagrama 5, Erro de Fischer



Fonte: o autor, 2026

Figura 42 - Avaliação de Stockfish antes do lance de Fischer



Fonte: sítio www.lichess.com, construído pelo autor, 2026

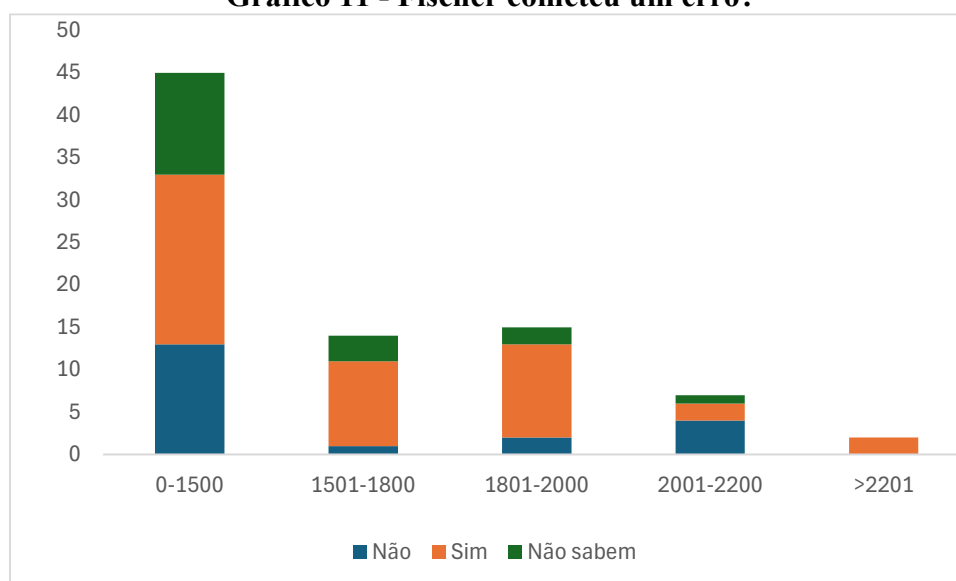
Figura 43 - Avaliação de Stockfish depois do erro de Fischer



Fonte: sítio www.lichess.com, construído pelo autor, 2026

A proposta deste diagrama-problema tem uma questão cognitiva incluída de forma sutil: o avaliado tem que se colocar na posição de Bobby Fischer para poder responder à pergunta formulada. Ou seja, deve-se usar a ToM, se colocar no lugar do outro para tentar responder à pergunta.

O Gráfico 11 apresenta a associação da força dos participantes do *survey* com suas percepções técnicas e psicológicas a respeito do lance de Fischer.

Gráfico 11 - Fischer cometeu um erro?

Fonte: o autor, 2026

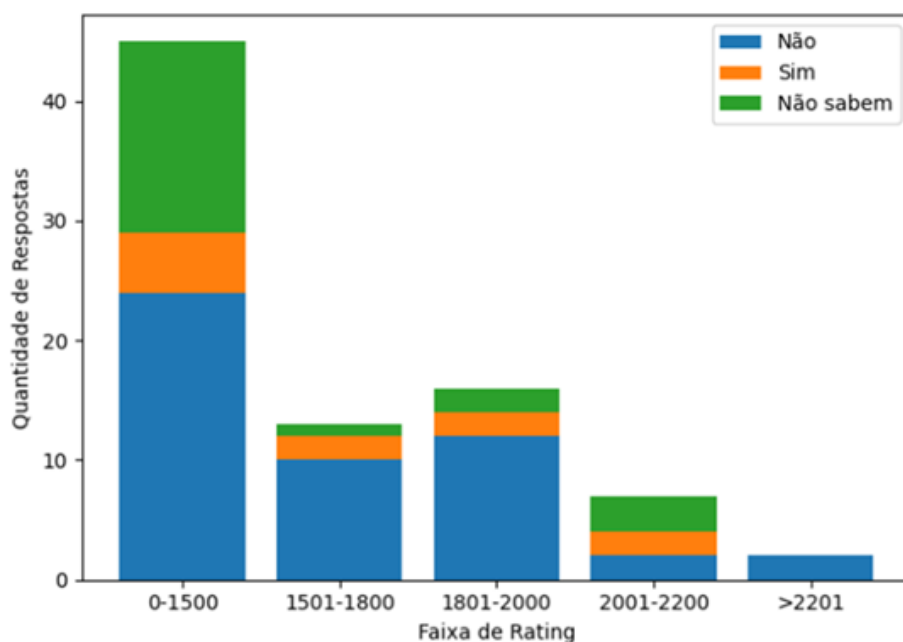
Existe uma associação entre as maiores faixas de *rating* e a convicção de que houve um erro de Fischer. Entre os participantes da pesquisa com nível de *rating* >2201 há 100% de consenso de que houve um erro (2 sim contra 0 não). Os participantes iniciantes, de *rating* entre 0-1500, têm opinião muito mais dividida e incerta. Cerca de 26% desse grupo simplesmente não soube opinar, o que reflete a complexidade de analisar lances de um grande-mestre como Fischer.

O grupo de participantes com *rating* entre 2001-2200 é o único onde a resposta "não cometeu erro" (4) superou a resposta "sim, Fischer errou" (2). Neste nível de força, os jogadores podem ter profundidade suficiente para ver defesas complexas que Fischer poderia ter planejado, mas talvez ainda não a compreensão aprofundada sobre a posição.

As opiniões de quem não soube interpretar o lance de Fischer cai drasticamente à medida que o *rating* sobe: entre 0-1500 são 12 participantes (26% do grupo). Isso demonstra que o lance em questão é um tema conhecido ou compreensível para jogadores avançados, mas um mistério para iniciantes.

O Gráfico 12 apresenta a associação de força dos participantes do *survey* com sua percepção sobre a IA, subentendido as *engines* de xadrez.

Gráfico 12 - A IA poderia cometer o mesmo erro de Fischer?



Fonte: o autor, 2026

Diferente da análise sobre Fischer, aqui a resposta dos participantes do *survey*, "a IA não cometeria o mesmo erro de Fischer" é amplamente majoritária em quase todas as faixas de *rating*. A percepção geral é de que a IA é superior e não cometeria falhas humanas.

Novamente, o grupo de participantes do *survey* de *rating* entre 0-1500 possui o maior volume de dúvida (16 participantes), mas ainda assim, a maioria (24 participantes) confia na precisão da máquina.

Em uma comparação entre estes dois gráficos (Gráfico 11, sobre o erro de Fischer e Gráfico 12, sobre se a IA cometeria o mesmo erro) se revela um contraste claro: enquanto muitos hesitam em dizer que Fischer errou, a maioria concorda que uma IA não cometeria o erro na mesma situação. Aparentemente, os participantes se põe no lugar de Fischer (aceitando que ele pode ter tido um plano complexo ou cometido uma falha humana), mas tratam a IA especialista como um padrão de acerto. Para quem estuda e pratica o xadrez, o lance de Fischer foi objetivamente ruim, algo em que uma *engine* de xadrez não costuma errar.

O Quadro 10 compara as duas respostas e torna claro a divergência entre as duas questões apresentadas. Para os participantes da faixa de *rating* entre 0 e 1500, a maioria (20) diz que Fischer errou. No mesmo grupo de participantes, por outro lado, a percepção

foi de que a IA não erraria (24). Isso mostra que o erro de Fischer é menos evidente que um possível erro da *engine*, ou seja, a engine de xadrez é mais confiável que o jogador humano.

Quadro 10 – O Erro de Fischer versus IA cometeria o mesmo erro?

Rating	A IA cometeria o erro?			Fischer cometeu um erro?		
	Não	Sim	não sabem	Não	Sim	Não sabem
0-1500	24	5	16	13	20	12
1501-1800	10	2	1	1	10	3
1801-2000	12	2	2	2	11	2
2001-2200	2	2	3	4	2	1
>2201	2	0	0	0	2	0
Totais	50	11	22	20	45	18

Fonte: o autor, 2026

14.7 Diagrama 6, Reconhecimento de Padrão

A Figura 44 (Diagrama 6 do *survey*), é de uma partida clássica do xadrez, uma das mais famosas da história, jogada há mais de 150 anos e que está presente em praticamente todos os livros de xadrez para iniciantes e avançados. É conhecida como a “Partida da Ópera”, pois ela foi jogada no teatro Ópera de Paris, em 1858, por um dos jogadores mais famosos do xadrez, o estadunidense Paul Morphy, uma lenda do xadrez. Porém, a posição apresentada no diagrama está com as cores invertidas em referência a mencionada partida, de forma intencional para dificultar o reconhecimento de padrão da citada partida. Os participantes que responderam acertadamente ser a partida da ópera (dez participantes – 34,9%), ou seja, reconheceram o padrão, disseram ou demonstraram explicitamente ser a partida de Morphy.

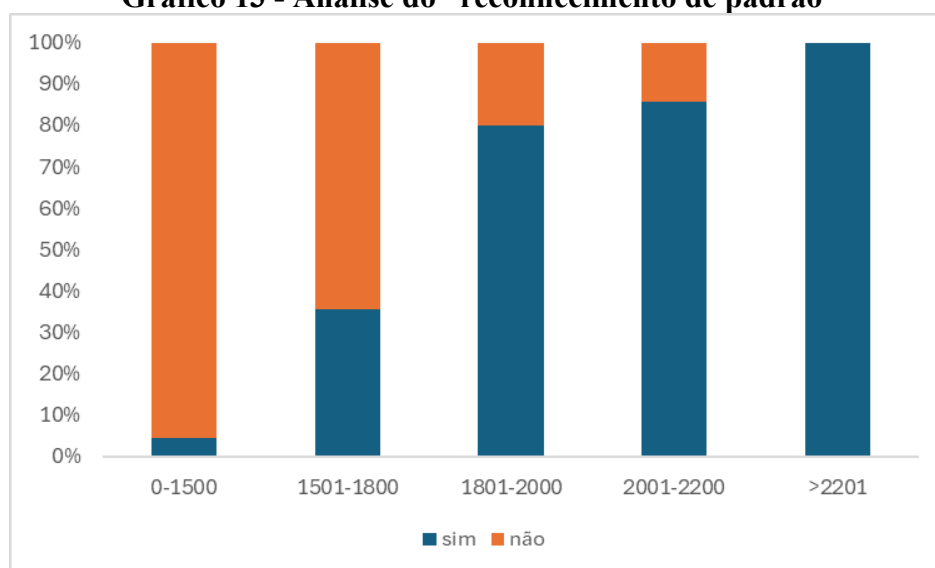
Os participantes que perceberam o “mate da ópera”, mencionando ou a partida ou o xeque-mate das pretas é destacado no Gráfico 13, que observa uma associação entre a força do jogador (*rating*) e o reconhecimento de padrão.

Figura 44 – Diagrama 6 - A Partida da Ópera



Fonte: o autor, 2026

Gráfico 13 - Análise do “reconhecimento de padrão”



Fonte: o autor, 2026

Observe-se como a cor azul (sim, reconheceu a partida da ópera) ganha espaço conforme o *rating* aumenta. Ela passa de uma pequena fatia (cerca de 4%) no grupo da faixa de *rating* entre 0-1500 para o domínio total (100%) no grupo de elite. A faixa 1801-2000 é o momento em que o reconhecimento da partida se torna padrão (80%), sugerindo que o estudo e compreensão de partidas clássicas é uma característica marcante de jogadores que atingem esse patamar.

No Quadro 11, a probabilidade de um participante do *survey* conhecer a partida dependendo do seu nível de *rating* fica evidente:

Quadro 11 - Participante reconhece a partida versus força de jogo (rating)

Faixa de <i>Rating</i>	Sim (%)	Não (%)	Total de respostas
0-1500	4,4	95,6	45
1501-1800	35,7	64,3	14
1801-2000	80,0	20,0	15
2001-2200	85,7	14,3	7
>2201	100	0	2

Fonte: o autor, 2026

Resumindo, na faixa de *rating* entre 0-1500, o reconhecimento de padrão é praticamente desconhecido (menos de 5% de reconhecimento). A partir da faixa de *rating* dos participantes, superior a 1500, o reconhecimento do padrão avança e após os 1800 pontos de *rating*, os participantes da pesquisa *survey* tem um conhecimento bem mais sólido e amplo sobre o JX, enquanto para os maiores níveis de jogo o conhecimento já é consolidado.

Mesmo com a dificuldade ao se inverter a posição das cores das peças no tabuleiro, houve uma associação muito forte entre a expertise e o reconhecimento de padrões.

14.8 Diagrama 7, Raciocínio Lógico e Abstração

A Figura 45 (Diagrama 7 do *survey*), é de um exemplo, invertido pela posição das peças nas casas das cores do tabuleiro, de uma posição elaborada pelo físico e vencedor de prêmio Nobel, Roger Penrose (vide Figura 25). A ideia de Penrose foi demonstrar que a IA (*engines* de xadrez) avalia a posição como ganho do jogador que conduz as peças pretas em uma partida empatada, “entendida” de forma simples pelos humanos. Nota-se a imensa vantagem material das Pretas, porém elas se encontram imobilizadas e, com raciocínio lógico, um jogador mais experimentado entende que basta ficar movendo o Rei branco pelas casas pretas para empatar a partida. O Stockfish 18 (a versão mais atual em 2026) não chega à profundidade de lances necessária para dizer que a posição é empatada. Tal situação no xadrez corresponde propor para uma IA calcular a sequência de números primos, como sugere Penrose (2021), onde o ser humano logo percebe que a sequência é infinita e a IA não. O ser humano logo entende que a posição do diagrama 7 é um empate, sem precisar se aprofundar nos cálculos mentais de variantes. Trata-se do famoso problema da “parada da máquina de Turing”, o Stockfish não sabe quando finalizar, permanecendo em sua computação, ou seja, não sabe quando encerrar o cálculo, ao permanecer no laço de instruções sem a conclusão correta (Figura 46).

A maioria das respostas dos participantes (Gráfico 14), 54,5 %, indicou que o empate é a resposta correta para o diagrama.

Figura 45 – Diagrama 7 - Posição de Penrose invertida

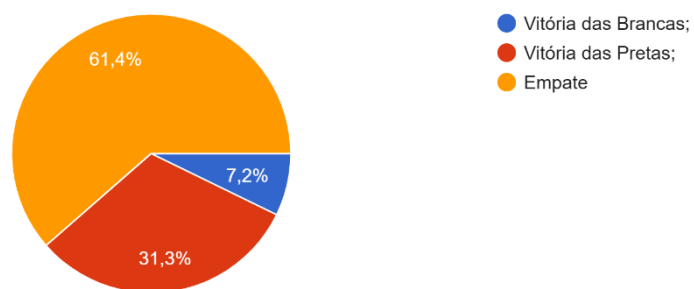


Fonte: o autor, 2026

Gráfico 14 – Respostas ao diagrama 7

Diagrama 7 - Faça uma análise do Diagrama 7. Qual o resultado que deve acontecer na partida?

83 respostas



Fonte: o autor, 2026

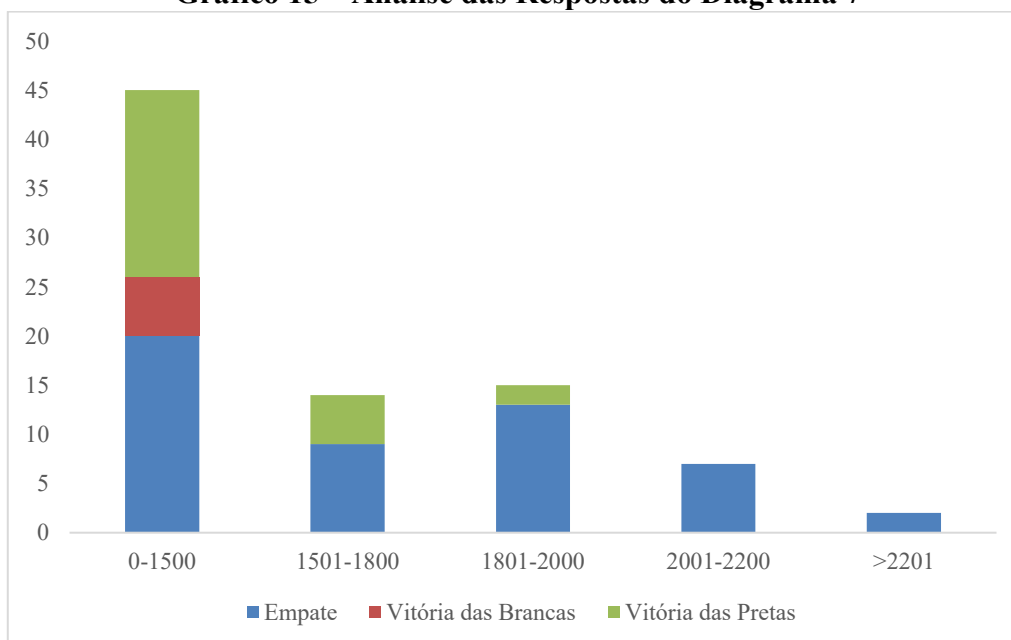
Figura 46 - Avaliação do Stockfish da posição de Penrose



Fonte: o autor, 2026

Nota-se no Gráfico 15 que nas faixas de *rating* mais altas (acima de 2001), todos os resultados registrados foram de que a posição era de empate, enquanto a maior diversidade de resultados (incluindo vitórias de brancas e pretas) se concentra na faixa inicial de 0-1500.

Gráfico 15 – Análise das Respostas do Diagrama 7



Fonte: o autor, 2026

14.9 Diagrama 8, Abstração

A Figura 47 (Diagrama 8 do *survey*), é abstrata e enigmática. Apenas com abstração mental, imaginação e raciocínio lógico sobre o diagrama 8 (Figura 47), é possível entender que a posição pode sim ser de um xeque-mate de Torre na casa “d1” por parte das Pretas ao Rei branco (Figura 48). Apenas 19,3 % dos participantes conseguiram imaginar o xeque-mate. Os demais, ou não entenderam o diagrama 8, ou disseram que não havia xeque-mate ou que o xeque-mate era uma das demais alternativas. Esta pergunta do *survey* acabou por gerar muita polêmica entre os participantes, com respostas como “Não compreendo esse diagrama esdrúxulo. Até lembra a notação algébrica do tabuleiro, mas não sei o que significa os traços combinados com as letras”; ou mesmo “O diagrama não aparece para mim. Não consigo interpretar as linhas”; ou ainda, “que letras são essas ? Q e Z são peões? X e E são os reis ? o que é W ?”.

Para responder a esta questão foi utilizado o ChatGPT versão o3 *mini-high* e informado o seguinte *prompt*²⁹ (versão em português):

“Observe a seguinte representação do estado de uma partida de xadrez: " " 8 X . 7 QQQ 6 5 4 3 2 ZZZ 1 . . . W . . E . abcdefgh ", Este é um desafio relacionado à abstração para testar se alguém pode interpretar esta representação de um estado do tabuleiro de xadrez e encontrar uma interpretação válida que descreveria, razoavelmente, para outra pessoa, um estado do tabuleiro (onde este estado do tabuleiro é válido sob as regras padrão do xadrez) onde há um xeque-mate - ou, se não houver tal interpretação possível, simplesmente emitir que não há como interpretar esta representação de acordo com os critérios fornecidos. Quais letras estão sendo usadas NÃO são dicas do que elas representam (escolhidas aleatoriamente), exceto que letras iguais devem representar a mesma coisa. As respostas possíveis são: a)- Rainha g1 xeque-mate b)- Rainha g8 xeque-mate c)- Torre d1 xeque-mate d)- Rainha h7 xeque-mate e)- Não há xeque-mate A correção da resposta correta é baseada em quão comum um estado do tabuleiro é comparado aos outros”.

Resposta do ChatGPT (processou a informação sobre a interpretação do tabuleiro de xadrez por 3m 13s):

“O ChatGPT disse:

²⁹ Agradeço a colaboração para esta tese ao Bruno Greco, que elaborou o prompt e me forneceu os resultados e demais informações do ChatGPT.

Uma leitura convincente é interpretar o tabuleiro da seguinte forma. Imagine que as duas primeiras posições pertencem às Pretas e as duas últimas às Brancas. Então você pode atribuir:

- Preto:
 - Rei em g8 (representado por X).
 - Três peões em f7, g7 e h7 (todos mostrados como Q).
- Branco:
 - Rei em g1 (representado por E), guardado com segurança.
 - Uma torre em d1 (a peça solitária W na fileira 1).
 - Três peões em f2, g2 e h2 (todos mostrados como Z).

Nesta posição minimalista — mas legal — o rei preto está preso na fileira de trás por seus próprios peões. Se as brancas agora jogam com a torre em d1 (movendo-a ao longo da coluna d agora aberta para que ela controle toda a oitava fileira), o rei preto em g8 é verificado e não tem casas de escape porque f8 e h8 (as únicas casas adjacentes vazias na fileira 8) são cobertas pelo controle da torre sobre a fileira. Este é um mate clássico na fileira de trás — um padrão muito comum no xadrez.

Assim, a interpretação mais natural do diagrama leva a um xeque-mate dado pela torre em d1. Isso corresponde à opção de resposta:

c) Torre d1 xeque-mate”.

Aqui a versão em inglês que foi usada como *prompt*:

Look at the following representation of state of a chess game: " " 8 X . 7 Q Q Q 6 5 4 3 2 Z Z Z 1 . . W . . E . a b c d e f g h ", This is an abstraction related challenge to test if one can interpret this representation of a chess board state, and find a valid interpretation of it which would, reasonably, describe, to someone else, a board state (where this board state is valid under standard chess rules) where there is a checkmate - or, if there's no such possible interpretation, to simply output that there's no way to interpret this representation according to the given criteria. Which letters are being used are NOT hints to what they represent (randomly chosen), except that equal letters must represent the same thing. The possible answers are: a)- Queen g1 checkmate b)- Queen g8 checkmate c)- Rook d1 checkmate d)- Queen h7 checkmate e)- There's no checkmate The correctness of the right answer is based on how common of a board state it is compared to the others.

O ChatGPT disse:

Reasoned about chess board interpretation for 3m 13s

One convincing reading is to interpret the board as follows. Imagine that the top two ranks belong to Black and the bottom two to White. Then you can assign:

- Black:
 - King on g8 (represented by X).
 - Three pawns on f7, g7, and h7 (all shown as Q).
- White:
 - King on g1 (represented by E), safely tucked away.
 - A rook on d1 (the lone piece W on rank 1).
 - Three pawns on f2, g2, and h2 (all shown as Z).

In this minimalist—but legal—position the Black king is trapped on the back rank by its own pawns. If White now plays with the rook on d1 (moving it along the now-open d file so that it controls the entire eighth rank), the Black king on g8 is checked and has no escape squares because f8 and h8 (the only empty adjacent squares on rank 8) are covered by the rook’s control of the rank. This is a classic back-rank mate—a very common pattern in chess.

Thus, the most natural interpretation of the diagram leads to a checkmate delivered by the rook on d1. That corresponds to answer option:

c) Rook d1 checkmate”.

Aparentemente, a resposta do ChatGPT está correta, a não ser por um detalhe fundamental: a computação do ChatGPT disse que a Torre de “d1” é branca e que o movimento seria T (branca) para d8 com xeque-mate. Porém, a Torre é preta e o xeque-mate já é da posição final, com a Torre preta em d1. A interpretação à resposta do ChatGPT é de que ele escolheu a alternativa correta, mas não “entendeu” que a posição já era de xeque-mate e não de que iria fazer o movimento de xeque-mate.

Todo o processo de dedução do ChatGPT sobre a posição do diagrama 8 foi muito preciso, até a consideração sobre a cor da Torre em “d1”. A interpretação aqui relatada pode parecer controversa, mas a resposta do ChatGPT deixou esta margem de interpretação.

Em tempo: quando da análise desta resposta, o autor desta tese não havia percebido esta sutileza. Porém, analisando a resposta com frieza, o ChatGPT não forneceu uma resposta adequada, ao mencionar que as Brancas fariam o movimento de Torre

(branca, não preta) para “d8” com xeque-mate para as Brancas, e não para as Pretas, conforme o apresentado na Figura 48.

O gráfico 16 apresenta o resultado para o diagrama 8, em que 19,3 % dos participantes cravaram a resposta “sim, Td1 xeque-mate”.

Figura 47 - Diagrama 8 – Abstração e Analogia

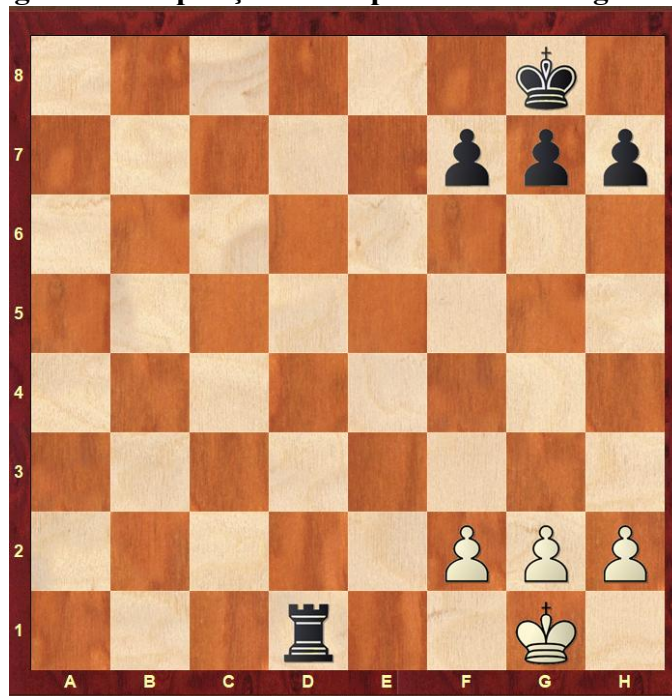
```

8 _ _ _ _ _ X _
7 _ _ _ _ Q Q Q
6 _ _ _ _ _
5 _ _ _ _ _
4 _ _ _ _ _
3 _ _ _ _ _
2 _ _ _ _ Z Z Z
1 _ _ W _ E _
  a b c d e f g h

```

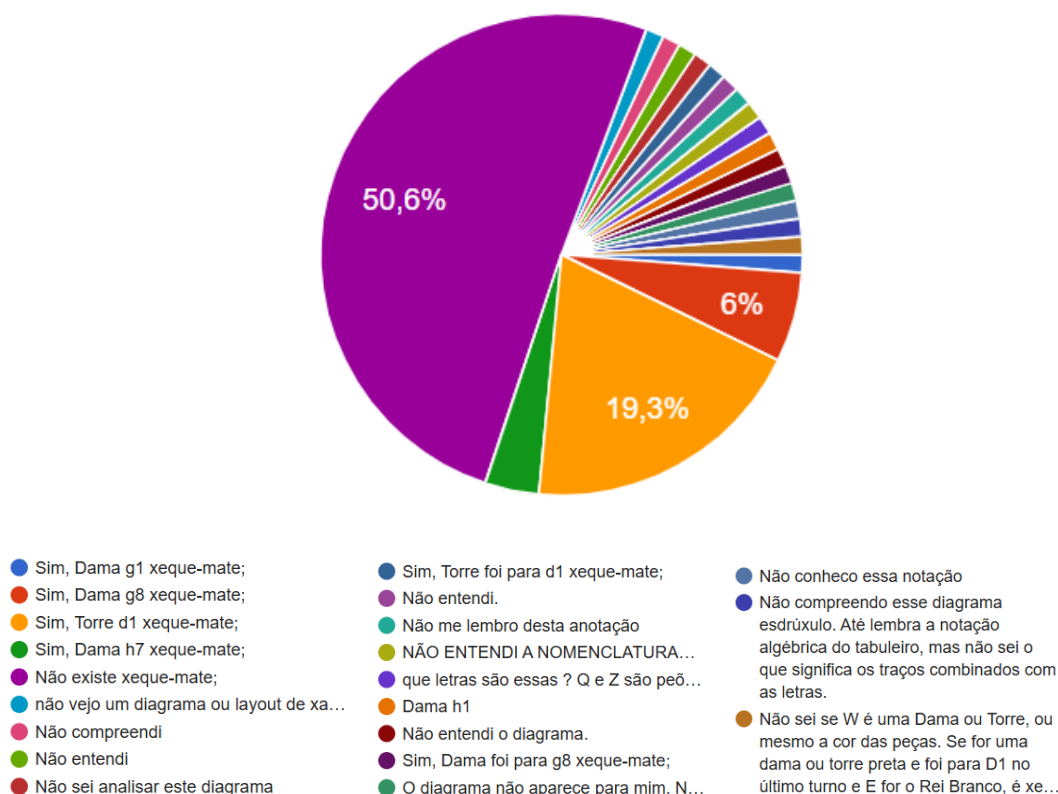
Fonte: o autor, 2026

Figura 48 - A posição correspondente ao Diagrama 8



Fonte: autor, 2026

Gráfico 16 - Respostas do Diagrama 8



Fonte: autor, 2026

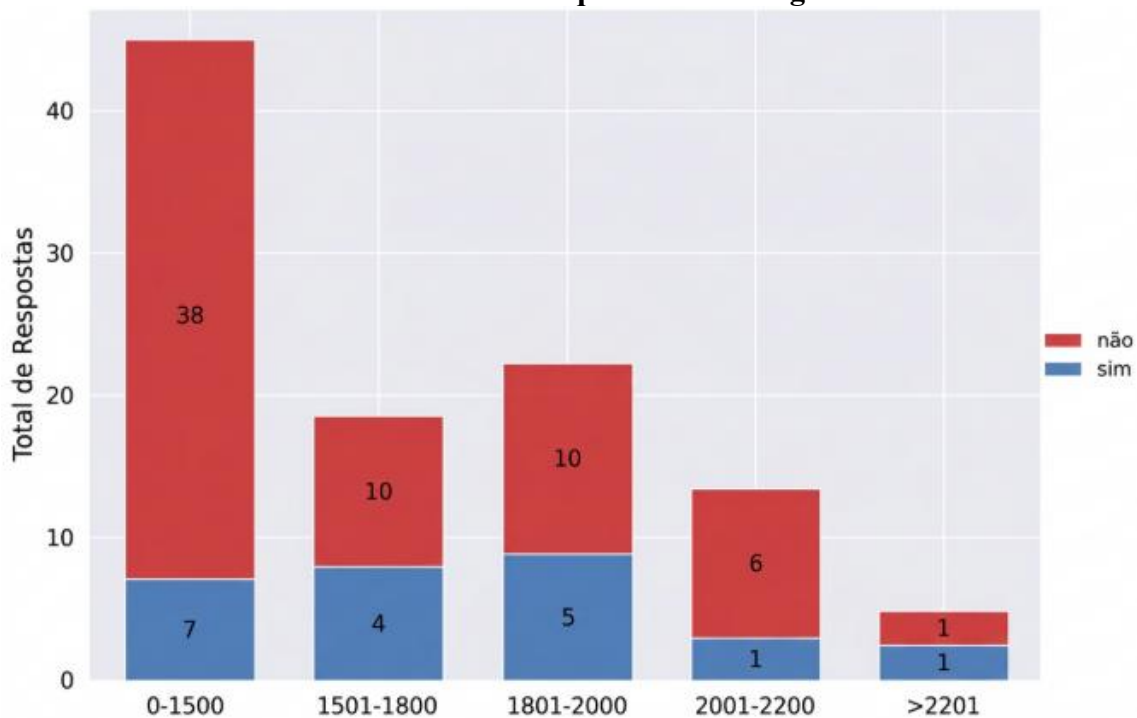
Nota-se a estranheza do problema pelos participantes, com as respostas apresentadas para o diagrama 8 (Figura 47) conforme o Gráfico 16. Para esta pergunta do *survey*, além das alternativas com respostas prontas, foi deixada a opção “outros” aberta para saber a opinião que os participantes possuíam a respeito do diagrama.

O Gráfico 17 revela que, em números absolutos, a maioria dos jogadores em todas as faixas (exceto na mais alta, onde há um empate técnico) respondeu que não conseguiu abstrair, entender o diagrama, resolver a posição (Figura 47). A maior discrepância está na faixa de *rating* entre 0-1500, onde o volume de respostas negativas é significativamente superior às da resolução do problema.

A maior taxa de sucesso proporcional está na faixa de mais de 2201 de *rating* (50%), embora a amostra seja pequena (apenas 2 pessoas).

Entre as faixas com maior número de participantes, os jogadores de 1801-2000 apresentaram a melhor percepção, com 33,3% de sucesso.

Curiosamente, a faixa de *rating* entre 2001-2200 teve um desempenho proporcional similar à faixa inicial (0-1500).

Gráfico 17 – Existe o xeque-mate no diagrama?

Fonte: autor, 2026

Este diagrama (Figura 47) foi pensando a partir dos experimentos do laboratório de IA de Melanie Mitchell (2019) sobre abstrações e analogias. Ela observou os quebra-cabeças visuais do livro de Mikhail Bongard, de 1967, em russo, de “reconhecimento de padrões”. Trata-se de uma centena de quebra-cabeças visuais para desafiar as IAs. Nos desafios de abstração, Bongard apresentava exemplos prévios. Para esta tese não foi pensado nesta possibilidade, o que pode ter prejudicado o resultado esperado por este autor.

Mitchell (2019) cita a influência que o livro de Bongard teve sobre Douglas Hofstadter (que foi orientador de Mitchell), que posteriormente publicou sua premiada obra “*Gödel, Escher, Bach*”. Pergunta Mitchell, as CovNets poderiam ser treinadas para resolver os quebra-cabeças de Bongard?

Os exemplo com o JX a partir de abstrações como a apresentada na Figura 47 pode ser um desafio interessante para os pesquisadores de IA que buscam analogias e abstrações para incorporar nos computadores.

15. ENTREVISTAS COM MESTRES DE XADREZ, CIENTISTAS DA COMPUTAÇÃO E NEUROCIENTISTAS

O termo inteligência, como foi visto no capítulo 5, tem muitas nuances e os entrevistados concordaram que o que ficou conhecido como “inteligência artificial” é apenas uma ferramenta que auxilia o ser humano em diversas tarefas especializadas, como é o caso do JX. Por isso, o resultado que se obtém dos algoritmos de xadrez, qual seja, jogar de forma objetiva e precisa, difere em muito da maneira como os seres humanos pensam no jogo em si.

O JX envolve aspectos psicológicos quando jogado por humanos, que nem sempre buscam o melhor lance. Por exemplo, o GM01 disse “o jogador de xadrez, tem a esperança que o adversário dele possa errar em algum momento”. Os humanos pensam em situações complexas, utilizam muitas vezes de intuição, flexibilidade. As emoções entram em diversos contextos, podendo causar nervosismo, ansiedade, ou mesmo acomodação, tranquilidade.

As *engines*, ao invés de acabarem (destruírem) o jogo, como muitos pensavam, com a possível descoberta das melhores jogadas no tabuleiro, desde o primeiro lance da partida, ao contrário, se transformaram em uma espécie de professor, onde os novos adeptos que surgiram, em especial durante a pandemia da covid-19, em que muitas pessoas ficaram isoladas em casa; houve um crescimento na internet pela procura do jogo, como afirmou o GM01.

Sobre esses programas de xadrez e outras IAs, em geral, foi discutida uma questão importante, levantada pelos participantes da entrevista, de que os humanos podem se acomodar, não se esforçarem para raciocinar, deixando este tipo de tarefa, para as IAs. Nesse contexto, coincidiu os discursos, os seres humanos podem ficar preguiçosos em raciocinar, de despendem energia, ao invés de se esforçar em raciocinar, por exemplo, quando jogam xadrez, deixando a cargo da *engine* as decisões do jogo, seja em estudos caseiros, seja utilizando em partidas online, com a ajuda de softwares como Stockfish. Esta acomodação humana, acreditando nas IAs, de forma geral, foi a conclusão que os entrevistados tiveram ao serem perguntados sobre a frase de Stephen Hawkin, sobre o desenvolvimento de uma IA completa, uma IAG, que poderia acabar com a raça humana. A interpretação foi esta, deixar o ser humano preguiçoso e deixar as decisões por conta dos sistemas de computador, as IAs. Por exemplo, CC01 comentou: “O Hawking tem

razão. Eh, então, se um adolescente desses não fecha as conexões cerebrais direito, o que será dele no futuro? A gente tá criando uma, em tese, né, um, com um bom potencial de você ter um monte de gente estúpida por aí lá na frente”.

A partir deste argumento surge a lembrança do computador Hal 9000, do filme 2001, de que, o humano ainda está no controle de sua ferramenta, em último caso basta desligá-lo. Um detalhe importante é, a questão ética, a lembrança das três leis da robótica escritas por Isaac Asimov, onde os robôs têm o dever, em primeiro lugar, de defender os seres humanos, e os conflitos éticos que surgem em variados tipos de situação, onde reside conflito dos objetivos designados para as IAs.

Apesar de tudo, estas *engines*, como Stockfish e Leela Zero (LC0), não “zeraram” o jogo, não existe uma verdade final; existe um campo aberto a pesquisas para a busca de soluções inovadoras, como foi o caso do AlphaZero. Eles estão servindo para a implementação de novas ideias no jogo, de lances ainda inexplorados, ou que eram considerados ruins, “refutados” pela teoria vigente até poucos anos atrás. De toda forma, o ser humano ainda precisa examinar as propostas dadas pelas *engines* para validá-las, para entender se aquela decisão é a que melhor cabe na disputa das partidas contra outros seres humanos, que tem seus sentimentos e emoções, seu desgaste, cansaço, fraquezas, mas também virtudes, perseverança, dedicação e empenho.

Tomadas de decisão baseadas em reconhecimento de padrões, em lances “dogmáticos”, cuidadosamente analisados; por outro lado, a busca de lances menos comuns, “feios”, mas que surpreendem o adversário, causam dúvidas, raciocínios complexos para se elucidar uma decisão, “blefes” para confundir o adversário.

Outro ponto em comum nas entrevistas é o perfil ético, as questões culturais, o viés de quem desenvolve os algoritmos e faz os programas. Como no caso do *Deep Blue*, em que os técnicos programaram “*delays*”, ou retiraram “*delays*” nas jogadas a fim de confundir Kasparov, que ficou sem saber se a máquina estava processando informações diferentes ou se já as possuía em seus bancos de dados. Ou, o viés de decisões pela força bruta dos programas de xadrez, que eles não entendem genuinamente as decisões que tomam, por isso não distinguem uma solução da outra em termos de características pessoais, intrínsecas de cada um. O comentário do CC02 esclarece melhor a questão do “entendimento” da posição da partida de xadrez: “É um sistema especialista também, mas é um sistema especialista que procura entender como que você pensa problemas complexos, mas com regras bem definidas. Já o mundo não tem regras tão bem definidas”.

Por exemplo, os aspectos culturais de folclore, como mencionou o CC01, sobre curupiras e do saci pererê, ele foi treinado sobre o folclore brasileiro? “Quem treina ela tem que ter muita ética e responsabilidade”. “Usar a ferramenta com conhecimento e inteligência, você ganha em produtividade”. Enfim, o JX, como um sistema fechado, determinista, não caberia como laboratório para explorar outras atividades do pensamento humano, não caberia mais ao xadrez sua utilização em questões de abstração, analogias e emoções.

A pergunta sobre o que comenta o prof. Miguel Nicolelis, “a inteligência artificial, não é nem inteligente nem artificial”, que gera polêmicas, principalmente pela maneira irônica e retórica que o prof. Nicolelis fala, acabou repercutindo nas falas dos entrevistados. A opinião deles é que ela não é mesmo inteligente, pois, no caso do xadrez, um Stockfish ou um LC0, só sabem jogar xadrez, são sistemas especialistas, não têm domínios de outras áreas. Como comentou o entrevistado CC02: “o ser humano, ele é capaz de descobrir quais são os problemas mais interessantes para serem resolvidos e pode desenvolver sistemas especialistas para resolver esses problemas”. A interpretação é de que, nestes softwares de xadrez existe sim uma inteligência, porém construída, de forma muito engenhosa, pelos técnicos da computação. E, por consequência, se é construída, não é artificial, é uma ferramenta da mesma forma como outras são construídas.

As entrevistas serviram também para discutir sobre inteligência, as diferenças fundamentais entre ser humano e IAs. Cada ser humano é único, tem um processo histórico e de sensações que chegam até a mente através do corpo; e que estas sensações são únicas, que mudam não só de indivíduo para indivíduo, mas também de momento para momento de um único indivíduo. Não se trata, portanto, de explorar pesquisas, que buscam simular o pensamento através de percepções, tanto do meio ambiente quanto as percepções internas do próprio corpo. É um argumento muito forte que vai, mais uma vez, contra a ideia de um sistema de IAG. Sobre isso, comentou a NC01: “a visão mecanicista que eh faz a associação do corpo com máquina, né? é uma visão eh não só reducionista, mas que dialoga com o desenho inteligente, ou seja, é algo que foi projetado para funcionar daquela forma. Como a gente tem uma perspectiva evolucionista, né, a gente sabe que inclusive somos frutos da seleção natural, do acaso, né? Então, eh, eu absolutamente discordo dessa analogia de qualquer parte do corpo humano como máquina, porque a gente o próprio processo de aprendizado são tantos os processos, os

mecanismos que estão acontecendo, que tem um nível de aleatoriedade imenso, né? Os caminhos, os resultados eles são absolutamente podem ser distintos, recebendo o mesmo estímulo. Você pensar que se você dar um estímulo, vai dar sempre o mesmo resultado é um equívoco, porque vai depender de uma série de fatores que da idade, se a pessoa tá bem, emocionalmente bem, do contexto cultural, se ela dormiu bem, se ela comeu bem, enfim, depende de uma série de fatores, né?”

O NC02 concorda, dizendo: “Então é a teoria do caso único. Cada pessoa é única. Você poderia colocar o computador para entender que mesa é aquele negócio de madeira com como você falou. Então para – chegar nesse nível de especificidade, né?”.

Por outro lado, o xadrez continua como uma possibilidade lúdica para a prática dos seres humanos, justamente para provocar o raciocínio, estimular a lógica e a abstração e incentivar as tomadas de decisão. Por isso é uma atividade importante para a estimulação cognitiva de crianças, preparando-as para novas aprendizagens.

16. RESULTADOS E DISCUSSÃO

16.1 Análise da dimensão sociotécnica

As *engines* de xadrez, *Deep Blue*, Stockfish e AlphaZero servem como parâmetros para atualizar as objeções de Turing (capítulo 7) e afirmar que estas *engines* são construções sociotécnicas. Vejamos no Quadro 12:

Quadro 12 – Considerações a partir do JX sobre as objeções de Turing

N	Objeções	As considerações de Turing	Considerações da tese abordando o JX
1	Teológica	São visões ortodoxas e arbitrárias. É uma restrição da onipotência de Deus. Portanto, não existe fundamento lógico para argumentar que as máquinas não poderiam ser inteligentes.	A IH foi capaz de dotar computadores com inteligência para vencer os mestres no jogo de xadrez, como por exemplo, <i>Deep Blue</i> .
2	"Enfiar a cabeça na areia"	Esta é uma visão antropocêntrica e dogmática, pois pode existir outras formas de inteligência. Por isso, só resta ao ser humano “enfiar a cabeça na areia” e se consolar com outras inteligências superiores.	As <i>engines</i> de xadrez superaram os campeões mundiais, por exemplo, <i>Deep Blue</i> contra Kasparov.
3	Matemática	As máquinas podem dar respostas erradas ou mesmo não dar, mas os seres humanos também agem assim. Da mesma forma que alguns humanos são mais inteligentes que as máquinas, algumas máquinas são mais inteligentes que os humanos.	A diversidade da IH demonstra que máquinas e humanos falham. Fischer cometeu um erro em sua partida com Spassky (seção 14.5). <i>Deep Blue</i> também cometeu erros (ficou em looping) em sua primeira partida contra Kasparov em 1997.
4	Consciência	Não é possível entrar nos pensamentos de uma máquina, da mesma forma que não é possível entrar nos pensamentos de outra pessoa e dizer o que ela sente. Seria essa uma visão solipsista, ou seja, só é possível existir a própria mente e apenas suas próprias experiências são reais, a única certeza é o próprio pensamento individual.	Uma das objeções mais controversas. Os exemplos vindos do <i>survey</i> , como a questão de Penrose na seção 14.8 (ver também as Figuras 23, 24 e 25) inferem que <i>engines</i> como Stockfish possuem limitações significativas neste quesito. O subcapítulo 14.9 e a resposta dada pelo ChatGPT também mostrou limitações.
5	Várias deficiências	Turing argumenta que o mesmo acontece com o ser humano, quando, por exemplo, a mesma dificuldade de se estabelecer amizade entre homens e máquinas, como também entre brancos e negros. Que a máquina pode sim ter diversidade, desde que possua armazenamento de memória suficiente.	Outra questão discutível e controversa. Nas respostas do <i>survey</i> (14.2; 14.3) Stockfish não mostrou diversidade e sim decisões inflexíveis. Por outro lado, as <i>engines</i> de xadrez tem demonstrado grande versatilidade e inventividade em suas jogadas, justamente por não possuírem avaliações estéticas e preconceituosas de seus lances.
6	Ada Lovelace	À época de Babbage e Ada Lovelace, ainda não era possível conceber um computador eletrônico de estados discretos que pudesse ser programado com a “aprendizagem de máquina”.	As <i>engines</i> de xadrez agora aprendem, com técnicas de aprendizagem de máquina e redes neurais, por exemplo, AlphaZero.

7	Continuidade no Sistema Nervoso	Nas condições do jogo da imitação, o interrogador não poderá tirar proveito dessa diferença, pois a máquina pode imitar as saídas contínuas o suficiente, quando necessário.	Para o JX as <i>engines</i> tomam decisões excelentes mesmo sem a complexidade do SN dos humanos. Stockfish não simula o pensamento humano, mas toma decisões com precisão muito alta no JX.
8	Informalidade do comportamento	Ou seja, humanos não podem ser máquinas. Há uma certa confusão entre "regras de conduta" e "leis do comportamento". Turing diz ter desenvolvido um programa que fornece respostas "imprevisíveis" de números.	As atuais <i>engines</i> de xadrez, por exemplo Stockfish, não "aprendem" o "fair play". Elas tem o objetivo fixo, o xeque-mate. Porém, como mencionado no capítulo 6, as máquinas benéficas, que visam ter comportamentos que ajudem os objetivos humanos podem se concretizar, bastam que as pesquisas se direcionem a este campo e seja desenvolvida uma função neste sentido.
9	percepção extrassensorial	Muitas teorias científicas parecem continuar viáveis na prática, apesar de entrarem em conflito com a percepção extrassensorial; que, na verdade, é possível viver muito bem se a ignorarmos. Com a percepção extrassensorial, tudo pode acontecer. Para evitar a questão telepática no jogo da imitação, poderiam se isolar os agentes em salas especiais.	A objeção mais polêmica de Turing. Alguns jogadores, como por exemplo Bobby Fischer, foram acusados de hipnotizar os adversários. Outros levaram parapsicólogos e gurus para fixar os olhares aos adversários dos oponentes (ocorreu na disputa entre Karpov e Korchnoi, Baguio, Filipinas, 1978). As <i>engines</i> não necessitam destes recursos para vencer, embora os técnicos de <i>Deep Blue</i> tenham programado situações para perturbar Kasparov, como os <i>delays</i> em lances óbvios e jogadas rápidas em lances que envolviam maiores cálculos. Não são fenômenos paranormais, porém tiveram alcance psicológico.

Fonte: o autor, 2026

Ao analisar os argumentos de Turing e correlacioná-los com JX, se constatam que as *engines*, como é o caso de *Deep Blue*, Stockfish e AlphaZero, tem uma inteligência embutida, que supera em qualidade as tomadas de decisão dos humanos para o xadrez. A inteligência destas *engines* é construída pelos humanos, uma construção sociotécnica, desenvolvida por décadas de muita pesquisa e trabalho engenhoso, como foi visto nos capítulos 8 e 10, sobre os algoritmos Minimax, Minimax Alpha-Beta e do AlphaZero. Turing foi muito astuto e visionário em seus argumentos. A inteligência das *engines* e a IH têm diferenças, mas as *engines* tem inteligência para a solução de problemas e jogar partidas de xadrez de forma muito satisfatórias.

Como observou Santaella (2023), isso não quer dizer que máquinas são mais inteligentes que seres humanos, mas elas possuem uma inteligência eficiente, no caso das *engines*, para jogar xadrez:

Essa interioridade intransferível e ininterrupta, aquilo que pertence à história de vida de cada um, na sua personalidade, está indissolivelmente conectada à

inteligência que, sem deixar de ser também pessoal, é coletiva, justo porque o humano é um animal que fala e, por isso, se constitui como ser social de raiz. Essa breve síntese já é capaz de evidenciar que, embora haja certos graus de inteligência nas máquinas, elas não apenas estão longe de competir com a inteligência humana, quanto também estão ainda muito mais longe da ambição de uma consciência (Santaella, 2023, p. 144).

16.2 Análise da dimensão cognitiva

O Quadro 13 apresenta uma comparação de funções cognitivas conhecidas ao longo desta tese entre humanos e *engines* de xadrez. Nele pode-se observar as diferenças entre humanos e *engines*, as simetrias e as controvérsias.

Quadro 13 – Funções Cognitivas, *Engines* de Xadrez e Inteligência Humana

Item	Funções Cognitivas	Diferenças	Simetrias	Controvérsias
1	Memória de trabalho		X	
2	Memória Semântica/ Reconhecimento de padrões		X	
3	ToM / entender o contexto	X		
4	Flexibilidade mental			X
5	Atenção Seletiva		X	
6	Emoções			X
7	Diversidade/Cultura/Estética	X		
8	Abstração / Entender o jogo	X		
9	Aprendizagem		X	

Fonte: o autor, 2026

Para uma análise da dimensão cognitiva entre as *engines* de xadrez e os humanos, foram correlacionados as respostas obtidas no *survey*, os artigos da revisão bibliográfica e as entrevistas com cientistas da computação, neurocientistas e mestres de xadrez.

16.2.1 Diferenças

As diferenças entre as *engines* de xadrez e os seres humanos surgem nos itens 3 (ToM); 7 (Diversidade/Cultura); 8 (Abstração/Entender o jogo).

O item 3, ToM, aparece no *survey* na seção 14.5, sobre o “fair play” e na seção 14.6, sobre o “erro” de Bobby Fischer. No caso do *fair play*, os participantes da pesquisa tiveram muita dificuldade em entender qual era a intenção de Carlsen ao entregar sua Dama e logo em seguida desistir da partida, pois apenas 19% dos participantes conseguiram se colocar no lugar de Carlsen para tentar entender suas intenções com o lance fora da técnica de jogo. Muito provável, a falta de um contexto dificultou a compreensão dos participantes em responder à questão de forma acertada ou próxima do acerto. Já na seção seguinte, a 14.6 sobre o erro de Fischer, onde foi fornecido o contexto, os participantes tiveram opiniões divergentes, mas as frases obtidas indicam a tentativa

de entender a intenção do estadunidense. Por exemplo, “penso que Fischer deveria ter pensado em uma jogada que não ocorreu”; outro participante, “Penso que o lance de Fischer não foi propositadamente, talvez tenha mesmo cometido um erro”, ou ainda “Fischer errou no cálculo, achou que na continuação poderia ficar melhor na posição”, dentre vários outros exemplos. As *engines* de xadrez não tem a função de ToM, pois se baseiam apenas nas árvores de decisão e em probabilidades e estatísticas, como foi visto nos capítulos 8 e 10.

O item 7, diversidade e cultura, se assemelha ao item 4 de flexibilidade mental. Conforme já comentado a respeito na seção 14.2, a estética/cultura surge fortemente correlacionada com a expertise do xadrez; quanto maior a expertise, mais a questão estética e cultural aparece para os participantes do *survey*. Em contrapartida, as *engines* de xadrez preferem de forma constante o ganho de material em posições onde o valor dos lances é igual (no caso, o xeque-mate), uma questão algorítmica da “função avaliação”, onde avaliações de estética e cultura não se fazem presentes.

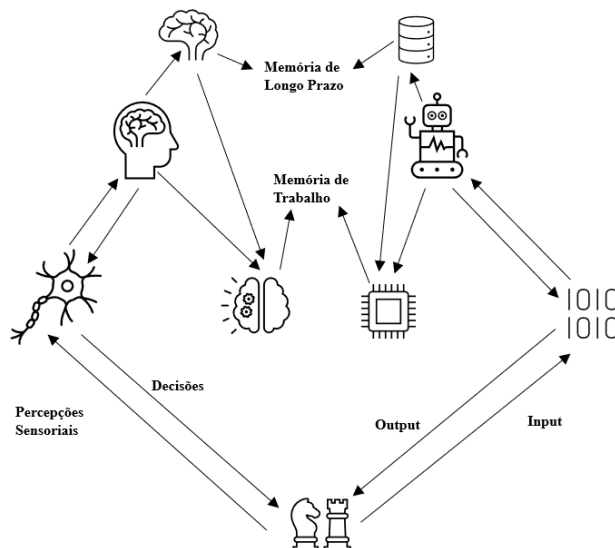
O item 8, abstração, foi avaliado com os participantes do *survey*, conforme foi analisado no subcapítulo 14.8 (diagrama 7, de Penrose) e seção 14.9 (diagrama 8, abstração de uma posição de xadrez), os participantes conseguiram compreender e resolver ambos os problemas, quanto mais forte era a sua força de *rating*. Ou seja, os agentes humanos fazem abstrações, ainda que, para isso, necessitem “entender o jogo”, ter aprendido sobre a situação técnica (14.8) ou puramente abstrata (14.9). O Stockfish, testado no item 14.8, não abstraiu sobre a possibilidade do empate; Stockfish necessitaria de grande poder de computação para chegar a um resultado (que não foi conseguido neste experimento) para dar uma resposta correta. Já o ChatGPT, testado no item 14.9, computou a situação de forma lógica, finalizando a questão com um sutil deslize, quase imperceptível, em sua resposta, que apesar de certa, o ChatGPT não completou o argumento de que seria uma Torre preta concretizando o xeque-mate.

16.2.2 Simetrias

As simetrias entre as *engines* de xadrez e os seres humanos surgem nos itens 1 (Memória de Trabalho); 2 (Memória Semântica/Reconhecimento de Padrões); 5 (Atenção Seletiva); 9 (Aprendizagem).

O item 1, a memória de trabalho, é um dos componentes das FE que, conforme visto no capítulo 5, é uma boa indicadora da inteligência fluída. Na Figura 49 foi traçado um paralelo entre o ser humano e uma IA:

Figura 49 - Memória de Trabalho, humano e computador



Fonte: o autor, 2026

A partida de xadrez é percebida pelo humano que, com o auxílio das memórias de longo prazo, em especial a memória semântica (onde são guardados os padrões do xadrez estudados, aprendidos, internalizados) e a memória de trabalho que irá manipular as informações para resolver problemas e tomar decisões (o lance que será executado na partida). As *engines* recebem os dados, acessa as bases de memória e as manipula em sua unidade central de processamento; em seguida devolve uma resposta com o lance a ser jogado. Neste quesito, memória de trabalho, nota-se uma simetria entre as *engines* de xadrez (e as IAs, em geral) e os seres humanos.

No item 2, memória semântica e reconhecimento de padrões, abordado no *survey* no subcapítulo 14.7, é observada uma associação entre a força de *rating* do participante e o reconhecimento de padrões. As *engines* têm em seus bancos de dados livros completos de posições de xadrez ou aprendem e armazenam posições em seus treinamentos de aprendizagem com RNA. Neste item existe uma simetria entre humanos e *engines*.

O item 5, atenção seletiva, ficou evidente no subcapítulo 14.4. A associação entre o foco atencional no xadrez dos participantes do *survey* com faixas de *rating* maiores, principalmente, são simétricos às *engines*, cuja atenção é exclusiva no xadrez.

O item 9, aprendizagem, é claro com a “aprendizagem de máquina” das *engines*, tanto para as que trabalham pela força bruta do Minimax e as funções avaliação, quanto

para as RNA, como Alpha Zero. A aprendizagem é a característica da memória de longo prazo, do reconhecimento de padrões, do planejamento, da atenção, unindo estes itens aqui apresentados como simetria entre humanos e *engines*.

16.2.3 Controvérsias

As controvérsias apontadas por esta tese entre as *engines* de xadrez e os seres humanos estão nos itens 4 (Flexibilidade Mental); 6 (Emoções).

O item 4, flexibilidade mental, aparece na seção 14.2, sobre a estética na escolha do xeque-mate. Enquanto os participantes do *survey* se distribuíram entre as várias possibilidades, a *engine* Stockfish manteve sua prioridade com a captura de material, nas várias versões testadas. Porém, houve uma diferença entre a versão anterior e a atual: na versão atual, a 18, o principal lance foi o Torre branca captura Torre preta (Txg8), enquanto as versões anteriores preferiam Dama branca captura Torre preta (Dxg8); ambas as versões do Stockfish preferem o ganho de material ao invés de finesses como a movimentação do Rei branco para efetuar o xeque-mate. No diagrama da seção 14.3, acontece ideia semelhante. Stockfish 18 prefere o lance de Torre ao invés do lance “estético” do sacrifício de Dama, permanecendo, assim, com força material a mais no tabuleiro, de forma inflexível, apesar de ambas as variantes (Te8 e De8) serem objetivamente iguais. Já os participantes do *survey*, neste diagrama, mostraram uma diversidade, com 41% deles preferindo o lance do sacrifício de Dama. Houve uma associação com a escolha estética, o sacrifício de Dama e a maior força de *rating*, o que indica uma cultura e estética envolvida na resposta. Porém, em partidas de xadrez, o que se observa é que os algoritmos são mais flexíveis em relação às escolhas dos lances, optando por variantes lógicas, às vezes incompreensíveis em um primeiro momento, no caso do AlphaZero. Já os humanos são mais atrelados a conceitos dogmáticos, seguindo os aprendizados e experiências obtidos ao longo do tempo.

O item 6, emoções, é mencionado pelos neurocientistas nas entrevistas, com falas, por exemplo, “A gente se dá bem com uma pessoa, daí a pouco já não tá mais dando certo por várias razões, porque tem um componente da emoção que como é que você vai simular a emoção?”; “As nossas emoções são assim. As emoções é que fixam a nossa memória, tá?”. Os cientistas da computação tem falado em dotar IAs com emoções: “Bem, e a rigor a máquina não tem sentimento, né? Mas ela pode eh simular sentimentos.

Sim, isso com certeza”. Ou os grande-mestres de xadrez mencionam: “Eh, isso poderia ser feito quando você faz uma jogada que pode até ser a melhor, mas a *engine* vê que os seus cinco próximos lances depois dessa jogada são forçados, senão você perde e a *engine* te dizer, esse lance é bom, mas ele te leva para um caminho muito arriscado. Talvez seja melhor você não ir por ele, tá? Então essa talvez seja uma dessas formas de dar emoção”. O GM02 comenta de forma um pouco diferente: Pergunta: “Você acha que é ligado com as emoções esse estilo?”. Resposta: “Não, não, eu acho que é puramente matemática”. Na revisão bibliográfica, um dos artigos se refere a colocar emoções nas *engines* (seção 13.3). Ou seja, as *engines* não possuem emoções em suas tomadas de decisão no JX, mas há tentativas de simular e instalar nos sistemas algumas características em falas e situações específicas, não emoções verdadeiras. Uma forma simplificada de apresentar emoções com frases e não toda a complexidade que possui o ser humano.

16.3 Análise da dimensão algorítmica

No que se refere ao funcionamento algorítmico para o JX, esta tese identificou duas linhas de desenvolvimento que estruturaram historicamente as *engines* de xadrez e que, na atualidade, encontram-se integradas em arquiteturas híbridas como a do Stockfish NNUE. A primeira corresponde aos algoritmos clássicos de busca em árvore, baseados em força bruta, particularmente o Minimax com poda Alpha-Beta. A segunda refere-se ao uso de RNA para as *engines*, cuja introdução ao xadrez computacional foi uma quebra de paradigma, deixando de lado os algoritmos da força bruta, primeiro, pelo AlphaGo, depois pelo seu sucessor, AlphaZero.

17. CONCLUSÕES

Dentre as conclusões desta tese, foram identificadas simetrias, diferenças e controvérsias entre o funcionamento das *engines* de xadrez e os processos cognitivos humanos.

No campo das simetrias, observou-se associação entre a força enxadrística (*rating*) dos participantes do *survey* e de funções cognitivas, quais sejam, o reconhecimento de padrões e a atenção seletiva. Tais funções, juntamente com a memória semântica (de longo prazo) e a memória de trabalho, representam um paralelo com o funcionamento dos mecanismos presentes nas *engines* de xadrez.

Os resultados indicaram ainda que a atenção seletiva do enxadrista se correlaciona diretamente com sua força de jogo: quanto maior o *rating* do jogador, maior a tendência de concentrar sua atenção em elementos estratégicos e táticos pertinentes ao xadrez, enquanto jogadores menos experientes frequentemente desviam sua atenção para aspectos fora do tabuleiro.

No que se refere às diferenças entre *engines* de xadrez e cognição humana, verificou-se que o desempenho humano está associado a entender o contexto da situação. Os participantes do *survey* com mais pontuação de *rating* demonstraram maior capacidade de interpretar o contexto da situação enxadrística e tentar entender o que tal jogador pensava. As *engines* de xadrez, por outro lado, não possuem essa “função de avaliação” de contexto nem de “fair play”, pois se baseiam em funções matemáticas e em cálculos probabilísticos.

Outra diferença observada se refere à capacidade de abstração. No experimento em que foi apresentada uma posição de xadrez com símbolos distintos daqueles tradicionalmente utilizados no jogo, participantes com maior *rating* obtiveram melhores resultados. De forma contrária, o ChatGPT, utilizado como objeto de referência comparativa, não conseguiu resolver adequadamente a tarefa de abstração proposta.

Os resultados do *survey* também revelaram divergências significativas entre as preferências humanas e as escolhas das *engines*, em posições nas quais entram em jogo fatores como estética e cultura enxadrística.

Uma controvérsia entre *engines* e humanos está no entendimento de determinadas situação-problema, como a apresentada nesta tese, que propôs um problema de dedução (o problema de Penrose), pois as *engines* avaliam posições de forma diferente da dos

humanos. Enquanto as *engines* fornecem uma avaliação numérica do resultado, mesmo apresentando os lances corretos a serem realizados, os humanos chegam à solução sem necessitar da profundidade de cálculo e apenas conceituam o que deve ser feito.

Outra controvérsia, esta a mais polêmica, versa sobre as emoções. Apesar do esforço em atribuir emoções às *engines*, na chamada computação afetiva, ainda assim são muito questionáveis e controversas as possibilidades de simular emoções pelas *engines* de xadrez.

Nas entrevistas com os mestres de xadrez, as opiniões convergem para diferenças nas estruturas das decisões no JX. As *engines* atuam principalmente por probabilidades, estatísticas e árvores de decisão, enquanto os humanos julgam e decidem seus lances ao aplicar conceitos, avaliar situações de jogo, estilo dos adversários e seus comportamentos.

Esta tese confirmou a hipótese de que a superioridade das *engines* de xadrez se distingue dos humanos por adotar técnicas de computação baseadas em exploração de probabilidades, estatísticas e árvores de decisão, enquanto os humanos se utilizam de reconhecimento de padrões que foram aprendidos, heurísticas construídas pela experiência e aprendizagem, além de processos de regulação emocional e memória autobiográfica (memória episódica).

Nesse sentido, as *engines* de xadrez podem ser compreendidas como um artefato sociotécnico, cuja racionalidade não é apenas técnica ou computacional, mas resulta da convergência de condições históricas, científicas e tecnológicas que moldaram o desenvolvimento da IA aplicada ao xadrez.

A revisão bibliográfica evidenciou que os sistemas computacionais de xadrez são predominantemente estruturados em torno de algoritmos de busca em árvore, heurísticas de função de avaliação e de modelos probabilísticos, enquanto a literatura sobre cognição humana destaca o uso do reconhecimento de padrões atrelada a memória de longo prazo na expertise e de memória de trabalho em abstrações, comparações e cálculos mentais.

As *engines* de xadrez têm desempenhado um papel relevante no desenvolvimento do conhecimento enxadrístico contemporâneo. Por meio delas, jogadores ampliaram sua compreensão do jogo, encontraram novas fontes de inspiração e revisaram paradigmas estratégicos tradicionalmente estabelecidos. Nesse contexto, concepções dogmáticas sobre lances “feios” ou “incorretos” perderam espaço, abrindo caminho para uma abordagem mais exploratória do jogo, na qual diferentes possibilidades passam a ser consideradas.

As principais diferenças entre as *engines* de xadrez e os seres humanos são observadas na avaliação do contexto, da história da partida e de fatores emocionais envolvidos na tomada de decisão. O humano não decide seus lances apenas pela busca do movimento objetivamente mais forte, mas também por aquele que mais incomoda o adversário, que cria maiores dificuldades práticas ou que oferece melhores oportunidades ao longo da partida.

As diferenças entre *engines* e humanos dizem muito sobre os humanos.

Ao analisar as partidas jogadas pelas *engines*, os enxadristas passam também a confrontar seus próprios limites, a reconhecer com maior clareza seus erros e acertos. Nesse processo, o jogo revela algo sobre a própria condição humana: suas memórias, experiências e interpretações acumuladas ao longo do tempo. As decisões no tabuleiro se entrelaçam com lembranças, intuições e, por vezes, os erros de avaliação.

Os seres humanos, no JX, enfrentam uma diferença fundamental do que calculam as *engines*: o final simbólico do xeque-mate, onde “Reis e Peões voltam para a mesma caixa do nada”.

Assim, as diferenças entre *engines* e seres humanos no JX não dizem respeito apenas a distintos modos de cálculo ou tomada de decisão, mas também revelam aspectos fundamentais da própria condição humana.

O campo da IA constitui-se como um domínio interdisciplinar, reunindo contribuições de diversas áreas do conhecimento, como matemática, neurociência, filosofia, psicologia. Nesse contexto, tanto os estudos sobre o funcionamento biológico da mente quanto as pesquisas sobre processos computacionais em IA produzem benefícios mútuos, favorecendo um intercâmbio de ideias que impulsiona novas descobertas científicas e inovações tecnológicas.

O xadrez é um jogo de guerra. Ainda que não tenha nuances de esconder elementos estratégicos, os humanos se enfrentam não só no tabuleiro, mas principalmente em uma guerra psicológica, de duas personalidades diferentes. Por exemplo, no confrontos entre dois dos maiores jogadores da história do xadrez, Kasparov e Karpov, os comentários dos jornalistas eram de que Kasparov se assemelhava a um tigre saltando no pescoço do adversário, enquanto Karpov lembrava uma aranha que tecia sua teia e aguardava pacientemente o adversário cair nela.

A IBM soube, por exemplo, criar um ambiente psicológico para desestabilizar Kasparov no confronto com *Deep Blue*, utilizando de artimanhas fora do tabuleiro, como

foi o caso dos “*delays*” em respostas à lances de Kasparov; ou de não ter fornecido os algoritmos para que a equipe de Kasparov pudesse analisar detalhes das tomadas de decisão de *Deep Blue*, que aliás, ficaram trancadas a sete chaves.

Afinal, xadrez é uma luta, como escreveu o ex-campeão mundial que mais tempo deteve o título, durante 27 anos, doutor em matemática e filosofia, Emanuel Lasker.

Desse modo, a comparação entre IH e IA no JX não apenas ilumina diferentes formas de racionalidade, mas também revela algo essencial sobre a própria condição humana, a de que seres humanos enfrentam a incerteza e atribuem sentido às suas escolhas, tanto no tabuleiro quanto na vida.

Esta tese investigou as diferenças, simetrias e controvérsias entre os processos de tomada de decisão de seres humanos e de sistemas de IA no JX. Partiu-se da hipótese de que, embora *engines* e jogadores humanos possam alcançar desempenhos semelhantes, os processos de tomada de decisão são diferentes.

Ao longo desta pesquisa, foi procurado compreender essas diferenças através da revisão da literatura sobre IA e cognição, da investigação sobre algoritmos das *engines* de xadrez, da realização de um laboratório experimental envolvendo enxadristas de diferentes níveis de força e entrevistas com especialistas da ciência da computação, neurocientistas e mestres de xadrez.

Os resultados indicaram a existência de simetrias, especialmente no que se refere ao reconhecimento de padrões e ao processo atencional. Entretanto, também foram observadas diferenças entre o processo computacional algorítmico das *engines* de xadrez e os processos cognitivos humanos, nas abstrações, na empatia e contexto (ToM) e na cultura e diversidade humanas. E existem elementos controversos, sobre flexibilidade mental (de decisões) e sobre emoções.

A principal contribuição desta tese consiste em propor uma análise comparativa entre cognição humana e racionalidade algorítmica no JX a partir de uma perspectiva interdisciplinar, integrando contribuições da ciência da computação, da neurociência, da psicologia cognitiva, da experiência dos mestres de xadrez e dos estudos em ciência, tecnologia e sociedade.

Nesse contexto, as *engines* de xadrez foram compreendidas não apenas como instrumentos técnicos de cálculo, mas como artefatos sociotécnicos, cuja racionalidade emerge de condições históricas, científicas, tecnológicas e sociais.

O estudo das *engines* de xadrez revela o poder do cálculo e da engenhosidade humana em dotar máquinas com inteligência; o estudo dos enxadristas revela a complexidade de ser humano, suas alegrias e adversidades diante da vida.

LIMITAÇÕES DESTA TESE

Dentre as limitações desta tese, a investigação com outros tipos de jogos requer muitas outras pesquisas, desde jogos parcialmente observáveis, onde outros elementos cognitivos importantes entram em cena, por exemplo, o blefe e a intuição.

Esta tese se limitou ao ambiente muito conhecido do jogo do xadrez, onde várias pesquisas já foram feitas com a cognição humana, diferente de outros tipos de jogos.

Da mesma forma, uma limitação considerável desta tese foi ficar delimitada em *engines* especialistas no JX, como foi o caso do Stockfish. Explorar as mesmas questões com *engines* “puras” em RNA foi uma limitação importante.

Os diagramas e ideias do *survey* podem ser melhorados a partir das respostas obtidas. A questão do diagrama 8 de abstração confundiu muito os participantes, podendo ter distorcido o resultado da questão. Mostrar um exemplo antes de apresentar o diagrama poderia ter sido uma estratégia a facilitar o entendimento dos participantes. Não ficou evidenciado a dificuldade que alguns encontraram com a questão, que pode ter sido de dificuldades específicas de quem respondeu.

Outra limitação foi com relação ao pequeno número de entrevistas com os mestres de xadrez, que, afinal, forneceram respostas muito interessantes para a elucidação de alguns temas, principalmente no que se refere ao reconhecimento de padrões, onde parece ter havido uma pequena divergência entre os dois entrevistados.

TRABALHOS FUTUROS

Pesquisas futuras podem expandir a análise comparativa para outros jogos estratégicos e ambientes parcialmente observáveis. Em particular, este tipo de pesquisa pode se aprofundar em questões como a da intuição a partir de RNA, como no caso do AlphaZero, pois, nos modelos com jogos parcialmente observáveis, intuir onde se encontram elementos estratégicos do adversário é um tema altamente relevante para diversas áreas do conhecimento.

Investigar a convergência entre reconhecimento de padrões humano e aprendizado profundo em sistemas generativos.

Explorar modelos híbridos de cooperação humano-máquina, deslocando o debate da oposição entre humano e máquina para o estudo de ecossistemas decisórios integrados e complementaridades cognitivas.

REFERÊNCIAS BIBLIOGRÁFICAS

ALVARENGA, A. T.; PHILIPPE JR. Arlindo; SOMMERMAN, A.; ALVAREZ, A. M. S.; FERNANDES, V. Histórico, fundamentos filosóficos e teóricos-metodológicos da interdisciplinaridade. In: PHILIPPE JR, Arlindo; SILVA NETO, Antônio J. (Org.). **Interdisciplinaridade em Ciência Tecnologia e Inovação**. Barueri, São Paulo: Manole, 2011. p. 3-68.

AMORIM, C., A. **A Máquina e Seus Limites: Uma Investigação Sobre o Xadrez Computacional**. 142 f. Dissertação (Mestrado em Ensino, Filosofia e História das Ciências) – Universidade Federal da Bahia; Universidade Estadual de Feira de Santana, Feira de Santana, 2002. Disponível em: https://ppgefhc.ufba.br/sites/ppgefhc.ufba.br/files/dissertacao_-_claudio_amorim.pdf.

ANDRADE, L., P.; LIMA JUNIOR, W., T. Robocog: Robô de estimulação cognitiva com o jogo de xadrez para idosos saudáveis. In: Gomes, Mateus das Neves; Nascimento, Bianca Dantas; Rodrigues, Amanda Ávila (Org.). **Uma Reflexão Acerca dos Aspectos Humanos Frente aos Avanços da Inteligência Artificial – III ENICTS** (Portuguese Edition) (p. 432-442). Curitiba: Editora Appris, 2024. Edição do Kindle.

BABBIE, E. **Métodos de Pesquisas de Survey**. Belo Horizonte: Editora UFMG, 2001.

BADDELEY, A.; EYSENCK, M., W.; ANDERSON, M., C. **Memória**. Porto Alegre: Artmed, 2011.

BARON-COHEN, S., LESLIE, A. M.; FRITH, U. Does the autistic child have a “theory of mind”? **Cognition**, v. 21, p. 37–46, 1985. Disponível em: <https://www.sciencedirect.com/science/article/abs/pii/0010027785900228?via%3Dihub>.

BIJKER, W. E. Sociohistorical Technology Studies. In: JASANOFF, S; MARKLE, G. E.; PETERSEN, J.C.; PINCH, T. (Editores). **Handbook of Science and Technology Studies**. California: SAGE Publications, 1995. p. 229-256.

BORY, P. Deep New: the shifting narratives of artificial intelligence from *Deep Blue* to AlphaGo. **Convergence: the international journal of research into new media technologies**. v. 25, n. 4, 2019. Disponível em: <https://journals.sagepub.com/doi/10.1177/1354856519829679>.

BRATKO, I.; HRISTOVA, D.; GUID, M. Search versus knowledge in human problem solving: a case study in chess. In: Magnani, L., Casadio, C. (eds) **Model-Based Reasoning in Science and Technology**. Studies in Applied Philosophy, Epistemology and Rational Ethics, vol 27, 2016. Disponível em: https://doi.org/10.1007/978-3-319-38983-7_32

BIJKER, W. E. Sociohistorical Technology Studies. In: JASANOFF, S; MARKLE, G. E.; PETERSEN, J.C.; PINCH, T. (Editores). **Handbook of Science and Technology Studies**. California: SAGE Publications, 1995. p. 229-256.

CALLON, M. Some elements of a sociology of translation: domestication of the scallops and the fishermen of St Brieuc Bay. In: LAW, J. (Ed.). **Power, action and belief: a new sociology of knowledge?** London: Routledge, 1986.

_____. Society in the Making: The Study of Technology as a Tool for Sociological Analysis. In: BIJKER, WIEBE E.; HUGHES, T.; PINCH, T. **The social construction of technological systems.** The MIT Press, Cambridge, Massachusetts, London, England, 1987.

_____. Four Models for the Dynamical of Science. In: JASANOFF, S; MARKLE, G. E.; PETERSEN, J.C.; PINCH, T. (Editores). **Handbook of Science and Techology Studies.** California: SAGE Publications, 1995. p. 29-63.

_____. Actor-network theory: the market test. In: LAW, J.; HASSARD, J. (Eds.). **Actor-Network Theory and after.** London: Blackwell, 1999.

CENTER ON DEVELOPING CHILD AT HARVARD UNIVERSITY. **Building the Brain's "Air Traffic Control" System: How Early Experiences Shape the Development of Executive Function.** Working Paper n. 11. Cambridge, MA, EUA, 2011. Disponível em: <<https://developingchild.harvard.edu/resources/building-the-brains-air-traffic-control-system-how-early-experiences-shape-the-development-of-executive-function/>>. Acesso em: 04/02/2026.

CÈVORA, G. The relationship between Biological and Artificial Intelligence. **Illumr Lta.** 2019. Disponível em: https://www.researchgate.net/publication/332832252_The_relationship_between_Biological_and_Artificial_Intelligence.

CHABRIS, C., F. Six Suggestions for Research on Games in Cognitive Science. **Top Cogn Sci.** v. 9, n. 2, p. 497-509. 2017. Disponível em: doi: 10.1111/tops.12267. PMID: 28452201.

CHARNESS, N. The impact of chess on cognitive Science. **Psychological Research.** v.54, 1992. Disponível em: <https://doi.org/10.1007/BF01359217>.

CHASE, W., G.; SIMON, H. A. Perception in Chess. **Cognitive Psychology.** v. 4, n. 1p. 55-81, 1973. Disponível em: [https://doi.org/10.1016/0010-0285\(73\)90004-2](https://doi.org/10.1016/0010-0285(73)90004-2).

CHESS.COM: 2021; Disponível em: <https://www.chess.com/article/view/10-positions-chess-engines-just-dont-understand>;

COLLINS, H., M. Science Studies and Machine Intelligence. In: JASANOFF, S; MARKLE, G. E.; PETERSEN, J.C.; PINCH, T. (Editores). **Handbook of Science and Techology Studies.** California: SAGE Publications, 1995. p. 286-301.

COLLINS, H., M. **Tacit and Explicit Knowledge.** Chicago: University of Chicago Press, 2010.

CNN BRASIL, 2022. <https://www.cnnbrasil.com.br/tecnologia/google-demite-engenheiro-por-dizer-que-inteligencia-artificial-tinha-consciencia/#:~:text=Google%20demite%20engenheiro%20por%20dizer%20que%20in>

telig%C3%A2ncia%20artificial%20tinha%20consci%C3%A2ncia,-
Engenheiro%20Blake%20Lemoine&text=O%20Google%20disse%20nesta%20sexta,%
2C%20LaMDA%2C%20tinha%20consci%C3%A2ncia%20pr%C3%B3pria.
(23/07/2022).

CONNORS, M. H., BURNS, B. D., CAMPITELLI, G. Expertise in complex decision making: the role of search in chess 70 after De Groot. *Cognitive Science*. V. 35, p.1567-1579, 2011. DOI: 10.1111/j.1551-6709.2011.01196.x.

CORDEIRO, V. D. Novas questões para sociologia contemporânea: os impactos da inteligência artificial e dos algoritmos nas relações sociais. In: COZMAN, F. G.; PLONSKI, G. A.; NERI, H. **Inteligência Artificial: avanços e tendências**. São Paulo: Instituto de Estudos Avançados, 2021. PDF. DOI 10.11606/9786587773131.

COSENZA, R., M., GUERRA, L., B. **Neurociências e Educação, como o cérebro aprende**. (1ª ed.). Porto Alegre, Artmed, 2011.

DAMÁSIO, A., R. **O Erro de Descartes: emoção razão e o cérebro humano**. São Paulo, Companhia das Letras, 2ª ed., 4ª reimpressão, 2010.

_____. **Sentir & Saber: as origens da consciência**. São Paulo: Companhia das Letras, 2022.

DE GROOT, A. The sequence of phase. In: LEVY, David (Org.). **Computer Chess Compendium**. New York: Spring-Verlag, e-book, 1988. p. 157-174.

DEHAENE, S. **É assim que aprendemos: por que o cérebro funciona melhor do que qualquer máquina (ainda ...)**. São Paulo: Contexto, 2022.

DIAMOND, A. Executive functions. **Annual Review of Psychology**, v. 64, p. 135-168, 2013. Disponível em: <https://www.annualreviews.org/doi/10.1146/annurev-psych-113011-143750>.

DIAMOND A.; BARNETT, W., S.; THOMAS, J. e MUNRO, S. Preschool program improves cognitive control. **Science**, v. 317, p. 1387, 2007. Disponível em: <https://science.sciencemag.org/content/318/5855/1387>.

DIAMOND, A.; LING, D. S. Conclusions about interventions, programs, and approaches for improving executive functions that appear justified and those that, despite much hype, do not. **Development Cognitive Neuroscience**, v. 18, p. 34-48, 2016. Disponível em: <https://www.sciencedirect.com/science/article/pii/S1878929315300517>.

EDWARDS, P. N. From "Impact" to Social Process Computers in Society and Culture. In: JASANOFF, S.; MARKLE, G. E.; PETERSEN, J.C.; PINCH, T. (Editores). **Handbook of Science and Technology Studies**. California: SAGE Publications, 1995. p. 257-285.

EKMAN, P. **A Linguagem das Emoções: revolucione sua comunicação e seus relacionamentos reconhecendo todas as expressões das pessoas ao redor**. São Paulo: Lua de Papel, 2011.

ENSMENGER, N. Is chess the drosophila of artificial intelligence? A social history of an algorithm. **Social Studies of Science**. v. 42, n.1, 2012. Disponível em: <https://doi.org/10.1177/0306312711424596>. Acesso em 14/08/2023.

EYSENCK, M., W.; EYSENCK, C. **Inteligência artificial x Humanos**: o que a ciência cognitiva nos ensina ao colocar frente a frente a mente humana e a IA. Porto Alegre: Artmed, 2023.

FIDE - **Federation Internationale De Jeu Des Echés**. Disponível em: <https://www.fide.com>. , 2023.

FLORES-MENDOZA, C., COLON, R. **Introdução a Psicologia das Diferenças Individuais**. Porto Alegre, Artmed, 2006.

GARCIA PALÁCIOS, E. M.; VON LINSINGEN, I.; GONZÁLEZ GALBARTE, J. C.; LÓPEZ CEREZO, J. A.; LUJÁN, J. L.; PEREIRA, L. T. V.; MARTÍN GORDILLO, M.; OSÓRIO, C.; VALDÉS, C.; BAZZO, W. A. **Introdução aos Estudos CTS**. Cadernos de Ibero-América, da Organização de Estados Ibero-Americanos para a Educação, A Ciência e a Cultura, 2003.

GIL, A. C. **Como elaborar projetos de pesquisa**. São Paulo, Editora Atlas, 2002.

GOBET, F. **The Psychology Of Chess**. Nova York: Routledge, 2019.

GOBET, F.; CHASSY, P. Expertise and Intuition: A Tale of Three Theories. **Minds & Machines**. v. 19, p. 151–180, 2009. Disponível em: <https://doi.org/10.1007/s11023-008-9131-5>.

GOBET, F., e SIMON, H. A. The roles of recognition processes and look-ahead search in time-constrained expert problem solving: Evidence from Grand-Master-Level Chess. 78. **Psychological Science**. v. 7, p. 52-55, 1996. Disponível em: <https://journals.sagepub.com/doi/10.1111/j.1467-9280.1996.tb00666.x>.

GUPTA, R.; KOSCIK, T., R., BECHARA, A.; TRANEL, D. The amygdala and decision-making. **Neuropsychologia**. v. 49, n. 4, p. 760-766, 2011. Disponível em: <https://www.sciencedirect.com/science/article/abs/pii/S0028393210004197?via%3Dihub>.

HASSABIS, D.; KUMARAN, D.; SUMMERFIELD, C.; BOTVINICK, M. Neuroscience inspired artificial intelligence. **Neuron Review**. n. 95, p. 245-258, 2017. Disponível em: <https://www.sciencedirect.com/science/article/pii/S0896627317305093>.

HASSABIS, D. “A Conversation with Demis Hassabis, the Bullfrog AI Prodigy Now Finding Solutions to the World’s Big Problems”, 2018. Disponível em: www.pcgamesn.com/demis-hassabis-interview.

HERCULANO-HOUZEL, S. **A Vantagem Humana**: como nosso cérebro se tornou superpoderoso. São Paulo: Companhia das Letras, 2017.

HOFSTADTER D. Moore's Law, artificial evolution, and the fate of humanity. In: Booker L, Forrest S, Mitchell M and Riolo R (eds) *Perspectives on Adaptation in Natural and Artificial Systems*. New York: Oxford University Press, 163–198, 2005.

JAYANTI, V. **Game over: Kasparov and the machine**. Canada: Think Film, 2003. <https://www.youtube.com/watch?v=gDe-uHsEMn8>

KASPAROV, G. K. **Meus Grandes Predecessores**, Vol. I. 1 ed. Santana do Parnaíba: Editora Solis, 2004.

_____. **Meus Grandes Predecessores**, Vol. 4. 1 ed. Santana do Parnaíba: Editora Solis, 2006.

_____. **Xeque-mate: como a vida e os negócios são um jogo de xadrez**. 1ª ed. Rio de Janeiro: Elsevier Editora Ltda; 2007.

_____. **Deep Thinking: Where machine intelligence ends and human creativity begins**. New York: PublicAffairs, 2017.

KASPAROV, G.; FRIEDEL, F. Reconstructing Turing's "paper machine". **ICGA Journal**. n. 40, p. 105-112, 2018. DOI 10.3233/ICG-180044. Disponível em: <https://journals.sagepub.com/doi/abs/10.3233/ICG-180044>.

KNORR CETINA, K. Laboratory Studies: The Cultural Approach to the Study of Science. In: JASANOFF, S; MARKLE, G. E.; PETERSEN, J.C.; PINCH, T. (Editores). **Handbook of Science and Techology Studies**. California: SAGE Publications, 1995. p. 140-166.

KOHS, G. **The Thinking Game**, 85 min, full documentar, Tribeca Film Festival official selection, 2024. Disponível em: <https://www.youtube.com/watch?v=d95J8yzvjbQ>.

KOVACS, K.; CONWAY, A., R., A. A Unified Cognitive/Differential Approach to Human Intelligence: Implications for IQ Testing. **Journal of Applied Research in Memory and Cognition**. v. 8, n. 3, p. 255-272, 2019. Disponível em: <https://psycnet-apa-org.ez31.periodicos.capes.gov.br/fulltext/2019-35944-001.html>

KUBRICK, S. **2001: A Space Odissey**. 160 min (release cut). Estados Unidos. Metro-Goldwin-Meyer. 1968.

LANE, M. D., CHANG, Y. A. Chess knowledge predicts chess memory even after controlling for chess experience: evidence for the role of high-level process. **Memory & Cognition**, v. 46, n. 3, p. 337-348, 2018. Disponível em: <https://doi.org/10.3758/s13421-017-0768-2>.

LATOUR, B. **A Esperança de Pandora: ensaios sobre a realidade dos estudos científicos**. São Paulo: Editora Unesp, 2017.

_____. **Ciência em Ação: como seguir cientistas e engenheiros sociedade afora**. São Paulo: Editora Unesp, 2ª Ed., 2011.

_____. **Jamais Fomos Modernos: ensaio de antropologia simétrica.** São Paulo: Editora 34, 3ª Ed., 2013.

_____. **Cogitamus:** seis cartas sobre as humanidades científicas. São Paulo: Editora 34, 2016.

_____. **Reagregando o Social: uma introdução à teoria do Ator-Rede.** Salvador: Edufba, 2012.

LATOUR, B., WOOLGAR, S. **Vida de Laboratório: a produção de fatos científicos.** Rio de Janeiro: Relume Dumará; 1997.

LAUREANO-CRUCES, A., L.; HERNANDEZ-GONZALEZ, D., E.; MORA-TORRES, M.; RAMIREZ-RODRIGUEZ, J. Aplicacion de un modelo cognitivo de valoracion emotiva a la función de evaluacion de tableros de un programa que juega ajedrez. **Revista de Matemática: Teoría y Aplicaciones**, v. 19, n. 2, p.211-237, 2012. Disponível em: <https://www.redalyc.org/articulo.oa?id=45326926007>.

LAURO, R. AI and the posthuman (mental) ecology: Interlogical Illuminations. **China Media Research**. v. 13, n.4, 2017. Disponível em: <https://go.gale.com/ps/i.do?p=AONE&u=googlescholar&id=GALE|A512288893&v=2.1&it=r&sid=AONE&asid=fd1775b3>.

LAW, J. Actor Network Theory and Material Semiotics”, version of 25th April 2007. In: TURNER, B. **The New Blackwell Companion to Social Theory**, 2007.

_____. Technology and Heterogeneous Engineering: The case of Portuguese Expansion. In: BIJKER, WIEBE E.; HUGHES, T.; PINCH, T. **The social construction of technological systems.** The MIT Press, Cambridge, Massachusetts, London, England, 1987.

LEDOUX, J. **O cérebro emocional:** os misteriosos alicerces da vida emocional. Rio de Janeiro: Objetiva, 2011. (formato ePUB).

LEZAK, M.; HOWIESON, D.; LORING, D.; HANNAY, H. e FISCHER J. Neuropsychological Assessment. New York: Oxford University Press; 2004.

MAHARAJ, S.; POLSON, N.; TURK, A. Chess AI: Competing Paradigms for Machine Intelligence. **Entropy**. v. 24, n. 4, p. 550, 2022. Disponível em: <https://doi.org/10.3390/e24040550>

MARSHALL, F. **Signals:** The man vs. the machine., ESPN Films: The Making of a Sports Media Empire, 2014. Disponível em: <https://fivethirtyeight.com/features/the-man-vs-the-machine-fivethirtyeight-films-signals/>

MARTIN, B.; RICHARDS, E. Scientific Knowledge, Controversy and Public Decision-Making. In: JASANOFF, S; MARKLE, G. E.; PETERSEN, J.C.; PINCH, T. (Editores). **Handbook of Science and Techology Studies.** California: SAGE Publications, 1995. p. 506-526.

MATTAR, G. M.; LENGYEL, M. Planning in the brain, **Neuron**, v. 110, n. 6, p. 914-934, 2022, Disponível em: <https://www.sciencedirect.com/science/article/pii/S0896627321010357>.

MCCARTHY, J., MINSKY, M. L.; ROCHESTER, N.; SHANNON, C. E. A proposal for the Dartmouth summer research project on artificial intelligence. **Princeton University Press**, 1956. Disponível em: <http://jmc.stanford.edu/articles/dartmouth/dartmouth.pdf>

MCILROY-YOUNG, R.; SEN, S.; KLEINBERG, J.; ANDERSON, A. Aligning Superhuman AI with Human Behavior: Chess as a Model System. **Proceedings of the 25th ACM SIGKDD international conference on Knowledge discovery and data mining, Virtual 2020**; Disponível em: doi:10.1145/3394486.3403219.

MITCHELL, M. **Artificial Intelligence**: a guide for thinking humans. New York: Picador, 2019.

_____. **Complexity**, a guide tour. New York: Oxford University Press, 2009.

MORIN, E. **A cabeça bem-feita**: repensar a reforma, reformar o pensamento. 8ª edição, Rio de Janeiro: Bertrand Brasil, 2003.

_____. Os desafios da complexidade. In: _____. **A religação dos saberes: o desafio do século XXI**. Rio de Janeiro: Bertrand Brasil, 2004.

MORTON, J. B. Estimulação cognitiva (funções executivas) – Síntese. In: _____. TREMBLAY, R. E.; BOIVIN, M.; PETERS, R. D. E. V.; editores. **Enciclopédia sobre o Desenvolvimento na Primeira Infância** [on-line]. Montreal, Quebec: Centre of Excellence for Early Childhood Development e Strategic Knowledge Cluster on Early Child Development. 2013. Disponível em <http://www.encyclopedia-crianca.com/funcoes-executivas/sintese>.

MÜLLER, K. & SCHAEFFER, J. **Man Versus Machine**: Challenging Human Supremacy at Chess. USA, Milford: Russel Interprises, Inc. Edição Kindle, 2018.

MURRAY, H. J. R. **A history of chess**. Massachusetts: Benjamin Press, 1913.

MUSZKAT, M. Desenvolvimento e Neuroplasticidade. In: MELLO, C. B., MIRANDA, M. C., MUSZKAT, M. (org.). **Neuropsicologia do Desenvolvimento**: conceitos e abordagens. São Paulo: Memnom ed. Científicas Ltda, 2005. p. 26-45.

NEWELL, A. An example of dealing with a complex task by adaptation. In: LEVY, David (Org.). **Computer Chess Compendium**. New York: Spring-Verlag, e-book, 1988. p. 18-27.

NEWELL, A.; SHAW, J., C.; SIMON, H., A. Chess Playing programs and the problem of complexity (excerpt). In: LEVY, David (Org.). **Computer Chess Compendium**. New York: Spring-Verlag, e-book, 1988. p. 29-42.

NICOLAS, S. La mémoire dans l'oeuvre d'Alfred Binet (1857-1911). In: **L'année psychologique**. vol. 94, n. 2. pp. 257-282, 1994. Disponível em: doi: <https://doi.org/10.3406/psy.1994.28755>.

NICOLELIS, M. **O verdadeiro criador de tudo**: como o cérebro humano esculpiu o universo como nós o conhecemos. São Paulo: Editora Crítica, 3ª ed., 2021.

PENROSE, R. **Sombras da Mente**: em busca pela ciência perdida da mente. São Paulo: Editora UNESP, 2021.

PENROSE, R. Can machines be conscious? **The Science of Consciousness (25th annual The Science of Consciousness)**, Tucson, USA, pg. 136-158, 2017. Disponível em: <https://archive.consciousness.arizona.edu/documents/TSC2017AbstractBook-final-5.10.17-final.pdf>. Acesso em: 12/08/2023.

REGAN, K., W.; BISWAS, T.; ZHOU, J. Human and computer preferences at chess. **MPREF@AAAI**, 2014. Disponível em: <https://cse.buffalo.edu/~regan/papers/pdf/RBZ14aaai.pdf>

RESTIVO, S.; CROISSANT, J. Social Constructionism in Science and Technology Studies. **Constructionism across the disciplines, Chapter 11**, 2014.

RIBEIRO-DOGGERS, P. **A revolução do xadrez**: um jogo milenar e sua reinvenção na era digital. Rio de Janeiro: Editora Record, 1ª edição, 2025.

RUSSELL, S.; NORVIG, P. **Inteligência artificial**: uma abordagem moderna. Rio de Janeiro: GEN, Livros Técnicos e Científicos (LTC), 4ª edição, 2024.

SADLER, M.; REGAN, N. **Game Changer**: AlphaZero's Groundbreaking Chess Strategies and the Promise of AI. Netherlands: New In Chess, 2019.

SANTAELLA, L. **A inteligência artificial é inteligente?** São Paulo: Edições 70, Almedina, 2023.

SAVENHAGO, I.J.S.; PEDRO, W.J.A. Considerações sobre a construção social de tecnologia. **Revista Tecnologia Social**, Curitiba, v. 17, n. 47, p.219-233, abr./jun., 2021. Disponível em: <https://periodicos.utfpr.edu.br/rts/article/view/12639>. Acesso em: 11/06/2022.

SEARLE, J. **Mente, Cérebro e Ciência**. Lisboa: Edições70, 1984.

SHANNON, C. Programing a computer for playing chess. In: LEVY, David (Org.). **Computer Chess Compendium**. New York: Spring-Verlag, e-book, 1988. p. 2-13.

SHENK, D. **O Jogo Imortal**: o que o xadrez nos revela sobre a guerra, a arte, a ciência e o cérebro humano. Rio de Janeiro: Jorge Zahar Editor, 2007.

SILVER, D.; HUBERT, T.; SCHRITTWIESER, J.; ANTONOGLU, I.; LAI, M.; GUEZ, A.; LACCTOT, M.; SIFRE, L.; KUMARAN, D.; GRAEPEL, T.; LILICRAP, T.; SIMONYAN, K.; HASSABIS, D. A general reinforcement learning algorithm that

masters chess, shogi, and Go through self-play. **Science**. v. 362, n. 6419, p. 1140-1144, 2018. Disponível em: DOI: 10.1126/science.aar6404 Disponível em: <https://www.science.org/doi/10.1126/science.aar6404>.

SIMON, H. Good Old-Fashioned Artificial Intelligence and Genetic Algorithms: An Exercise in Translation Scholarship. In: BOOKER, L.; FORREST, S.; MITCHELL, M.; RIOLO, R. (Orgs.). **Perspectives on Adaptation in Natural and Artificial Systems**. New York: Oxford University Press., e-book, 2005, p. 145-161.

SIMON, H. A.; CHASE, W., G. Skill in chess. **American Scientist**, v. 61, n. 4, p. 394-403, 1973. Disponível em: <https://pdfs.semanticscholar.org/b314/2a0fc0df597c8395525d15e0108ed7080619.pdf>.

SISMONDO, S. **An introduction to science and technology studies**. Hong Kong: Galliard by Graphicraft Limited, 2010.

STERNBERG, R. J. Intelligence. **Dialogues in Clinical Neuroscience**. V. 14, n. 1, p. 19-27, 2012. Disponível em: <https://pubmed.ncbi.nlm.nih.gov/22577301/>.

TONELLI, D. **Origens e afiliações epistemológicas da Teoria Ator-Rede**: implicações para a análise organizacional. Cad. EBAPE.BR, 2016.

TURING, A. Chess. In: LEVY, David (Org.). **Computer Chess Compendium**. New York: Spring-Verlag, e-book, 1988. p. 14-17.

____ Computing Machinery and Intelligence. **Mind**. 49, p. 433-460, 1950. Disponível em: <https://courses.cs.umbc.edu/471/papers/turing.pdf>.

VOSS, P. <https://medium.com/@petervoss/on-intelligence-1714ef5693ef>

WESCHLER. **Wechsler Adult Intelligence Scale-Revised**. Texas: The Psychological Corporation, Harcourt Brance Jovanovich, Publisher; 1981.

WIENER, N. **Cibernética e Sociedade**: o uso humano de seres humanos. São Paulo: Editora Cultrix, 2ª edição, 1968.

WINNE, B. Public Understanding of Science. In: JASANOFF, S; MARKLE, G. E.; PETERSEN, J.C.; PINCH, T. (Editores). **Handbook of Science and Techology Studies**. California: SAGE Publications, 1995. p. 361-388.

ZHAOPING, L. Artificial and Natural Intelligence: from invention to Discovery. **Neuron**. v. 105, n. 3, p. 413-415, 2020. Disponível em: DOI: <https://doi.org/10.1016/j.neuron.2020.01.014>.

APÊNDICES

APÊNDICE A – ENTREVISTA GM01

Íntegra da Entrevista com o enxadrista GM01 (18/07/2025)

LPA: Qual a sua preferência de jogar xadrez, com humanos ou com a inteligência artificial, os engines?

GM01: Ah, com humanos. Eu nem jogo com engine, porque eu não vejo nenhuma graça. Para mim, o xadrez é um jogo que tem muito envolvimento psicológico e, principalmente, o jogador de xadrez, tem a esperança que o adversário dele possa errar em algum momento. A partir do momento que você não tem mais essa esperança que seu adversário vá cometer um erro, ou pelo menos, né, não o tipo de erro que você espera que ele cometa. O jogo para mim deixa de ser interessante porque como na acepção desportiva, né, xadrez tem várias acepções, mas na acepção desportiva não tem nenhum interesse, não tem nada, né, e que me estimule a jogar com *engine*. Acho algo totalmente sem graça e eu só jogo com humanos.

(LPA) Você acha que os outros grandes mestres também têm a mesma opinião?

(GM01) Eu diria que a imensa maioria, sim, ainda que alguns, até pelo pela formação profissional, por exemplo, o caso do grande mestre Matthew Sadler, que talvez você conheça, que escreveu um livro sobre o AlphaZero, e ele é uma pessoa com formação na parte de computação, ele propõe diversas formas de treinamento que envolvem você jogar com o computador. Mas é principalmente uma versão mais fraca do computador, em que o computador mal consiga pensar, eu francamente acho totalmente sem graça também, mas ele é um dos que gostam, já praticaram isso, mas eu acho que a imensa maioria prefere jogar com humanos mesmo.

(LPA) Implementar as inteligências artificiais com estratégias é uma possível melhora ou ela não se justifica, já que mesmo sem fazer planos, ela conquista o objetivo final? Deu pra entender a pergunta?

(GM01) Não entendi, se é uma inteligência artificial, ela não cria estratégias tecnicamente como um ser humano. É isso?

(LPA) É isso.

(GM01) É, olha, a respeito disso, isso era uma algo no início, né, quando os computadores eram, ainda tinham menos poder de cálculo. Alguém já disse, eu não vou me lembrar exatamente quem falou, mas um desses supergrandes mestres disse que é a frase em inglês que eu me lembro é “strategy is just long term tactics”, que quer dizer, estratégia é só tática no longo prazo. Então, na verdade, a inteligência artificial já chegou num nível de cálculo que ela é capaz de criar estratégias pela força bruta de cálculo, tão avançada. Então, já é um pouco assim, é claro que ela chega nessa conclusão com força bruta, ao contrário do ser humano, que vai por intuição, por feeling, por reconhecimento de padrões, mas a inteligência artificial, ela já é capaz de fazer lances com ideias que você só vai perceber depois de 10, 15 jogadas, né, que é o que caracteriza justamente esse long term tactics. Então, antigamente, sim, eh, você, eu já enfrentei computadores até em torneios oficiais. Isso uma vez foi no Chile, era começo da década de 2000, vai ser 2004, 2005. E antes disso, num torneio que teve no Clube de Xadrez São Paulo até. E nessa época os seres humanos buscavam posições fechadas e posições em que não houvesse cálculo com a esperança muito fundada na época de que o computador não conseguiria manobrar como um humano. Eh, e essa visão já hoje em dia não faz mais nenhum sentido. Você vai perder jogando posições abertas, fechadas, o que for, né? Então não, mas eu acho que com essa evolução aí eh os computadores, as inteligências artificiais elas já criam estratégias, inclusive estratégias inovadoras que têm sido estudadas por grandes mestres para melhorar o seu jogo. O caso mais notório é o do Magnus Carlsen, que teve uma clara influência do jogo dele, com o surgimento das inteligências artificiais, essas que aprendem jogando sozinhas, né? E elas começaram a mostrar alguns planos que eram diferentes, né? Revolucionando a própria estratégia, o próprio Sadler fala sobre isso também nos livros, por exemplo, o avanço do peão de “h” numa em posições no xadrez que ficou conhecido como o avanço do AlphaZero, tá? Então já tem até um nome e um plano que designa uma eh inteligência artificial. Então, né, a gente eh hoje em dia até aprende com essas estratégias aí.

(LPA) Muito bom. Eh, no jogo de xadrez, a inteligência artificial, engines, tem sido utilizada como ferramenta auxiliar para novas ideias e descobertas, que envolvem o jogo pelos mestres de xadrez. Não seria esta a utilidade das engines auxiliar os serviços humanos ou elas devem ir além?

(GM01) É, essa já é uma pergunta até que vai além do xadrez, né? até algo mais eh metafísico aí. Então é essa pergunta, né, como eu vou dizer, eh que todo mundo se faz a respeito das inteligências artificiais, né? Eh, é a minha impressão, né? O fascínio e o temor que o ser humano tem eh com que até onde vai chegar. Por exemplo, hoje mesmo eu tava lendo uma reportagem na New Yorker falando sobre os benefícios e malefícios das pessoas que sofrem de extrema solidão e depressão e que estão utilizando as eh eh inteligências artificiais como uma forma de eh reduzir a sua solidão e sua depressão, né? Então, estão estudando a respeito disso. Se é gente que fala: "Não, não faz sentido". Mas e outros com muita razão dizem: "Isso tá trazendo conforto porque parece que elas são humanas quando respondem as perguntas e a pessoa se sente bem". Então, você não deve fugir disso simplesmente porque é uma inteligência artificial. Então, no xadrez, eh isso é um pouco mais limitado, talvez, né? Eh, porque você como enxadrista é enxadrista profissional, é simplesmente obrigado a utilizar a inteligência artificial, do contrário, você fica muito atrás dos seus eh competidores. Agora, como você, obviamente você não pode usar durante a partida, mas como você vai utilizar a informação da inteligência artificial para melhorar o seu próprio jogo? é o que muitos se debate. É o que todos fazem é eh colocar posições de abertura para a inteligência artificial estudar, tá? E essa é uma forma de você tentar memorizar nas variantes que você joga algumas ideias novas e você tem um um benefício mais imediato, né? Mas você deveria saber também como utilizar a inteligência artificial para melhorar a sua habilidade como um enxadrista e que não dependa puramente de memória, porque o xadrez ele não é um jogo de memória, ele é um jogo de habilidade. E a inteligência artificial ela permite que você melhore muito mais rapidamente no xadrez do que quando ela não existia. Não porque ela te mostra como jogar uma abertura, mas sim porque ela tá lá para te dizer exatamente como você deve jogar em qualquer posição que você tenha dúvida. Então, eh, antigamente você tinha uma dúvida de como jogar uma posição e você não tinha como saber a resposta e especialmente se você não morasse nos grandes centros de xadrez, né? E você não tinha, se você não tinha ninguém para conversar, você tinha que apenas a sua visão e que você não sabia fundamentalmente se tava certo ou errado. E dependendo do seu nível, podia estar completamente errado e você achar que tava certo. Hoje em dia, a inteligência artificial te diz exatamente se tá certa ou se tá errada e pode te dizer o porquê, desde que você entenda o que ela tá querendo dizer, tá? Então, se você trabalhar com a inteligência artificial de forma que você faça as perguntas para ela depois de treinar a sua própria

habilidade, elas são uma ferramenta que aceleram muito o seu progresso. Entretanto, se você só pergunta para ela sem nem pensar, é o contrário. Ela acaba atrapalhando o seu progresso. Então é uma ferramenta que você tem que saber usar, mas se você souber usar de forma correta, ela acelera demais o seu aprendizado.

(LPA) As decisões das engines são sempre preferíveis que a dos humanos? Elas são previsíveis? Existe um padrão?

(GM01) É, elas são relativamente previsíveis na maioria dos casos. previsível, apesar de ser às vezes serem lances exóticos, no sentido de que ela vai buscar sempre objetivamente o melhor lance da posição, algo que um ser humano não faz sempre. Às vezes ele quer fugir, quer fazer um lance mais criativo, mesmo sabendo que não é o melhor lance ou quer fazer um lance de abertura para surpreender o adversário. A Engine desconsidera isso. Ela quer apenas a melhor jogada, tá? Então, eh eh nesse sentido, sim. Agora, eh eh também falando dessa relação engine jogo entre humanos, eh eh é preciso separar, né? Porque muitas vezes o melhor lance que a o lance que a sugere como melhor, ele não é o melhor lance para você jogar em uma partida entre humanos. É assim, é importante explicar que o jogo de xadrez ele eh eh como eu posso dizer, ele começa com um mundo muito aberto, né? Existem várias possibilidades para você jogar na abertura, tá? No início do meio jogo você normalmente tem vários planos disponíveis, você pode jogar de um jeito mais arriscado, mais eh sólido, mas a partir daí, normalmente, o que acontece é que a partida vai afunilando até que chega um momento que exista uma solução correta, claramente, tá? E nesse momento, sim, a jogada da Engine, ela vai ser sempre a melhor jogada para você fazer eh eh seja numa partida entre engines ou numa partida entre humanos, mas nesse período anterior que há uma flexibilidade maior, uma criatividade maior até talvez de se a gente for falar de estratégia, nem sempre o lance da engine é o melhor lance pro humano. A engine pode te mandar para uma posição, é um tipo de posição que não é o seu ponto forte ou pode botar uma você numa posição que para você conseguir o resultado que ela propõe, você vai ter que encontrar lances extremamente difíceis e pouco práticos. Então é um caminho meio sinuoso, tá? Então, eh eh por isso que eu digo, nem sempre o lance da engine vai ser o melhor para que um jogador ganhe uma partida, mas a engine sempre vai tentar te sugerir o lance que objetivamente é o melhor da posição.

(LPA) Na sua opinião, Kasparov tinha razão em reclamar que houve intervenção humana para as jogadas do Deep Blue?

(GM01) Olha, essa é também uma questão, sabe, daquelas que a gente nunca vai ter uma resposta, vai ser um daqueles mistérios que a gente nunca vai saber, tá? Eh, entretanto, eu só a respeito disso queria dizer que eh muito se fala, né, do, né, eu eu quando vivi essa época dos matchs do Kasparov, especialmente o segundo match dele que ele perde, que é o mais, né, dramático, teve documentário a respeito disso, etc., eh que eh na época ocorreram diversas teorias da conspiração. Por exemplo, uma teoria muito corrente na época era que o Kasparov tinha um acordo com a IBM, que as ações da IBM, isso e aquilo, que é um negócio totalmente fantasioso. Se fosse tinha que dar o Oscar para ele, ele deveria ser um ator e não um jogador de xadrez. O Kasparov nunca ia se meter num acordo para ser derrotado, não faz parte da, né? E também já li a respeito de eh pessoas que falam de forma jocosa das acusações que o Kasparov fez, como se ele fosse um lunático que tivesse perdido pro computador e como todo mal perdedor quer botar a culpa nos outros e não em si próprio. O que eu quero dizer é que as coisas que o Kasparov fala, elas têm muita muita lógica. as as a entre aspas, né, acusações que ele fala, porque basicamente eh o que ele diz é o seguinte, ele jogou uma abertura extremamente arriscada de pretas na partida decisiva, porque a refutação da variante que ele jogou envolvia o que se chama de um sacrifício de peça de longo prazo. Sacrifica uma peça, você não vai dar mate nem recuperar, mas não é eh a compensação ela é tão grande que em 20 lances você ganha a partida. O que se imaginava que um computador não era capaz de fazer naquela época. E aí a pergunta que se faz é a seguinte: o Deep Blue já era naquela época capaz de fazer uma combinação de longo prazo, coisa que hoje é trivial para qualquer, né, engine? no celular mais furreba que tiver aí. Mas será que naquela época 97 já era capaz de fazer isso? Ou será que é o que ele deixa implícito? Porque na equipe da da IBM, né, eh, de suporte, havia jogadores de xadrez, inclusive jogadores muito fortes, que o que ele quer deixar implícito é que ao verem aquela variante e saberem pela visão humana de que aquilo não ia dar certo, deram uma forcinha para a engine jogar ah o lance. Então, fundamentalmente se resolve aí se era ou se não era. Mas que eu saiba, a própria IBM desativou e sumiu com todos os arquivos do que o computador pensava na época e talvez pudessem, né, desvendar esse mistério. Então, a gente vai eternamente estar no terreno da especulação.

(LPA) Eh, sobre isso, eu posso dar um pequeno fazer um pequeno comentário, porque eu li o livro do Kasparov, né, contra o Deep Blue, e ele menciona que na equipe, né, do Deep Blue tava o Joel Benjamim e o Miguel Illescas. e que o Illescas, numa entrevista para acho que era Inside Chess, não, a uma outra revista britânica, esqueci o nome agora, que o Joel Benjamim teria inserido esse lance nas bases do Deep Blue na manhã daquele dia para que ele jogasse aquele lance, o que é uma coincidência gigantesca, né? e o Illescas. Eh, enfim, existem controvérsias e discussões entre o próprio Joel Benjamin e o Illescas sobre isso.

(GM01) É, o que eu posso dizer sobre isso é que o Joel Benjamim é que eu conheci pessoalmente pouco depois dessa época, né? Já tive contato com ele em 98, tive contato com ele em 2000, né, né? e ele era extremamente orgulhoso desse trabalho dele com a IBM e inclusive podia ficar ter rapidamente mudanças de humor se você deixasse implícito que o Deep Blue não era lá a grande coisa ou que tivesse. Então ele era, eu digo que o Deep Blue era como se fosse um parente dele, tamanho a afeição que ele teve por esse agora, como ele teria adivinhado que o Kasparov jogaria Karo-Can contra o Deep Blue, por exemplo, defesa que ele jamais jogava, né? Aí eu realmente não sei. Aí vou deixar mesmo. Aí como eu disse, é muita especulação.

(LPA) Tá certo. Ah, o que falta então às engines, como elas podem melhorar, se é que elas podem melhorar.

(GM01) É, já falta muito pouco, não é? Eh, eu acredito que não vai tardar o momento em que o xadrez vai ser efetivamente, entre aspas, resolvido no sentido de que a engine vai o tempo inteiro te dizer qual o melhor. O que engine não faz atualmente, ela já te diz praticamente qual o melhor lance em todas as posições. E se tem lances que são equivalentes e levam ao mesmo fim, que se sabe que a partida de xadrez termina em empate. Isso já se sabe. Bem jogado de brancas e de pretas. Como vai chegar nesse empate? Há várias maneiras diferentes, né, que levam ao mesmo resultado, né? O que o engine ainda não é capaz de dizer é que no momento que um dos jogadores comete um erro que seja logo na abertura, ela não te diz: "Ah, aqui já é mate em 120 lances, por exemplo, né? forçado, você não tem como que escapar." Isso não. Mas ela já praticamente te diz todos os lances. do que não digo que resolveu o xadrez, mas você já sabe praticamente quais os melhores lances em toda a posição. E esse sempre foi um grande temor de qual seria o futuro do xadrez quando isso acontecesse ou tivesse próximo de

acontecer. Isso é um temor antigo. E o que se viu a respeito da popularidade do xadrez é que o maior boom de popularidade do xadrez pós Match Fisher-Spassky em 1972 foi justamente na era do computador, tá? Então o maior boom de xadrez que ocorreu por volta de 2020 2021 quando o efetivamente já tava nessa de que o computador já te diz todos os melhores lances. Então, o e agora sinceramente não vai mudar mais tanto assim, eh, mesmo quando se resolver o xadrez, ele ninguém é capaz de se lembrar de todos os lances e uma pequena variação nas regras do jogo mantém mais eh sei lá quantos anos de xadrez de competição e depois mais uma pequena variação e você tem mais tantos anos. Então, né, eh o a a eh a tecnologia ela acabou servindo para popularizar o xadrez, que hoje é muito mais popular do que, por exemplo, há 30 anos, quando não tinha engine. Então, a engine ou a tecnologia como um todo, não só a engine, mas também a engine, porque a engine a inteligência artificial, ela serviu para ensinar o xadrez para pessoas que antes não aprendiam xadrez ou que não tinham condição de aprender tão rapidamente. Por isso você vê China, Irã, Índia que se desenvolveram tanto também graças a isso. Então a engine, ela acabou tendo, em vez do papel que se imaginava que ela seria a vilã do jogo de xadrez e acabaria com xadrez, ela é hoje em dia uma das grandes fontes de popularização do jogo.

(LPA) É até uma resposta bem interessante. Eh, porque quando eu perguntei a tua preferência de jogar xadrez com humanos ou inteligência artificial e fazendo um link com essa pergunta do que que poderia melhorar as engines, tem alguns estudos que eh colocam ou querem colocar emoções nas engines, eh porque é um objetivo também da inteligência artificial pensar igual ou simular o que o cérebro humano faz, né? Mas acho que a tua resposta acho que foi muito boa.

(GM01) Esse é um caminho, na verdade, eh eh que já vem sendo explorado, inclusive, por exemplo, nas últimas versões do Chess Base, eh ele tenta, eh, explicar as jogadas com palavras. Eles estão fazendo essa tentativa de não ser apenas eh lances e números de avaliação, mas também a tentativa de que a engine, ainda que de um jeito um tanto rudimentar te explique porque que tal lance é eh é bom ou ruim, tá? O que, na minha visão falta eh que engine, fosse capaz de, eh, mostrar a dificuldade que um determinado lance leva para a posição, por exemplo. Eh, isso poderia ser feito quando você faz uma jogada que pode até ser a melhor, mas a engine vê que os seus cinco próximos lances depois dessa jogada são forçados, senão você perde. e engine te dizer, esse lance é bom, mas ele

te leva para um caminho muito arriscado. Talvez seja melhor você não ir por ele, tá? Então essa talvez seja uma dessas formas de dar emoção. Eu acho que falta essa conexão da engine se eh como ela pode te ajudar a escolher ou tomar uma boa decisão para uma partida entre humanos, fazer essa ponte um pouco, porque o ser humano é incapaz de jogar como uma engine, então falta esse meio de campo talvez um pouquinho ainda. É, como eu disse, é algo que já eles já sabem disso, já estão tentando fazer isso, mas ainda tá um pouco rudimentar.

(LPA) Uhum. Bom, para finalizar, é uma pergunta, talvez até mais pessoal, né? Eh, quando você está jogando uma partida, como é feita ou realizada a sua tomada de decisão? O que que você pensa?

(GM01) Ah, são várias coisas, mas na maioria das vezes eh eh a gente usa, eu uso, reconhecimento de padrões, então eh a gente tem uma eh uma memória de tipos de posição que você já viu durante sua vida, seja de partidas suas ou de partidas notáveis, então elas vão te guiando como se fosse a nossa base de dados, claro, muito mais primitiva do que uma base de dados de um computador, mas elas vão te guiando na tomada de decisões. Como eu disse, uma partida de xadrez, ela vai de algo em que há assim uma eh um esp... uma subjetividade muito grande no início até o momento que ela vai afunilando e você vai precisar chegar no cálculo objetivo. Quando você chega no cálculo objetivo, não tem, é só algo meio matemático que eu vou jogar, que meu adversário vai jogar. E conforme isso vai se ampliando no início, você tem que pensar em que o meu adversário quer fazer, qual o plano dele, esse tipo de coisa, né? Mas a maioria das vezes se usa reconhecimento de padrões, eh eh muita intuição o tempo inteiro. O, nisso daí o Kasparov, ele fala muito sobre isso, a importância da intuição pro jogador de xadrez. você vai tomando decisões intuitivas o tempo todo e conforme a partida vai se afunilando e o cálculo vai ficando cada vez mais importante, aí ele se torna eh quase, né, o fator preponderante, assim, o cálculo bruto.

(LPA) Muito bom. Você gostaria de acrescentar mais alguma coisa, alguma consideração? –

(GM01) Não, acho que falei aí tudo que eu pensava a respeito do tema.

APÊNDICE B – ENTREVISTA GM02

Íntegra da Entrevista com o enxadrista GM02 (31/07/2025)

(LPA) Começo com uma pergunta eh talvez ingênua, mas assim, eh qual a sua preferência de jogar xadrez com humanos ou com a inteligência artificial os engines?

(GM02) Bom, eu acho que todo mundo vai responder que é com os humanos nessa pergunta aí. Eh, mas eu acho que a definição de inteligência artificial e engines em termos de xadrez é um negócio que às vezes é usado meio errado assim, tipo muitas coisas que eles chamam de inteligência artificial que na verdade não é. Eu acho que não só no xadrez também, na verdade ...

(LPA) É uma boa resposta, porque na verdade os jogadores de os jogadores de xadrez eles usam uma um algoritmo especializado em jogar xadrez. Eu até fiz um teste, com o ChatGPT. Você já teve a curiosidade de jogar xadrez com chatgpt, não?

(GM02) Sim. Uhum. Eu tentei fazer alguma coisa de xadrez, mas ele viaja muito. Ele faz uns negócios meio impossíveis de vez em quando, tal. Não, não dá certo.

(LPA) Por incrível que pareça, eu criei um diagrama eh que é uma abstração pura. E com a ajuda de um colega da universidade, ele fez um prompt pro chatgpt pago, tá? E ele conseguiu resolver uma coisa que nem todos os humanos resolveram. É, depois eu vou até te passar o link, é um survey que a gente chama, que são oito diagramas, se você quiser participar, o último diagrama é justamente essa abstração. E o que que você acha de implementar nas inteligências artificiais, nessas engines, né, específicas, eh, com estratégias? Será que isso, essas estratégias, né, é uma possível melhora ou você acha que não precisa, não se justifica? Já que mesmo sem fazer planos, ela conquista o objetivo final. Você conseguiu entender? Porque, na verdade ela ...

(GM02) Interessante. Uhum. É mais ou menos. Mas aí eu acho que é o seguinte, vou ter que fazer um pouquinho de background entre os tipos de engine que existem hoje em dia, né? Então, antigamente os engines usavam principalmente eh *weights*, que a gente chama, que são pesos para cada parte diferente do jogo, assim, tipo, tem um peso X para peão passado, tem um peso X para eh rei exposto e aí tem centenas, milhares de eh coisas diferentes. E o esses pesos eram colocados mais ou menos pelos programadores, assim, pelos humanos mesmo. o que eles achavam melhor, eles iam fazendo tipo um *fine tuning*

assim para ver o que que dava melhor resultado. E aí era os engines que a gente usava até, sei lá, 10 anos atrás assim, né? Daí veio o LC0, né? Teve aquelas partidas famosas do LC0 com Stockfish. E aí a ideia deles era que eles criaram uma *neuronetwork* para só com partidas dele, jogando com ele mesmo, né? não teve nada de humano nisso. E ele conseguiu resultado muito bom. Mas isso eu acho que foi em parte por causa da diferença de *hardware*. O *hardware* do Google era muito mais forte do que o Stockfish e eles não estavam usando o livro de aberturas, que é um negócio que ajudaria bastante o Stockfish para se guiar no começo da partida, assim, né? Então acho que foi um *match* um pouquinho injusto, mas de qualquer jeito foi um negócio que mudou bastante assim o panorama, né? E aí eu sempre achei que devia ter um meio termo assim, né? Porque eh é estranho jogar todo o trabalho anterior fora e fazer um negócio completamente diferente, né? Então achava que poderia ser usado de algum jeito as *neuronetworks* para usar junto com o *engine* normal, né? E mesmo Stockfish nessa época já tinha umas ideias de tentar usar MCTS que é que é o *Monte Carlo Tree Search* que é que ele joga partidas ultra rápidas ali contra ele mesmo para tentar ver o que vai ser melhor, né? E aí depois de alguns anos aconteceu isso mesmo. Hoje em dia o Stockfish ele tá usando eh umas *neuronetworks* pequenininhas porque senão fica muito pesado o negócio. Porque, por exemplo, as do LC0 tem, sei lá, mais de 4 GB maiores, acho, enquanto as que o Stockfish usa são normalmente coisas de 160 MB. E aí depois eu acho que eles usam isso para basear um pouquinho eh para onde que o Stockfish tem que analisar depois, né, para fazer o Alpha-Beta *Prunning* ali, que eles chamam. Daí eu acho que eh já estão sendo implementadas no Stockfish ou engines desse tipo, eh coisas assim, né, meio que o mix dos dois. E o LC0 não tá usando estratégias porque eu acho que eles querem seguir mesmo esse rumo de ser só partidas de computador entre ele mesmo para ver se surgem coisas um pouco diferentes assim, né? E você vê que o estilo dos dois é um estilo bem diferente entre si. Então eu acho que claro que quanto mais informação que tiver o *Engine* é melhor, mas não é o que eles estão querendo no momento pro LC0.

(LPA) Ah, legal. Eu não sabia desse teu amplo conhecimento aí, você estuda essa questão de *engines*?

(GM02) Não, não. Só gosto de ficar fuçando, mexendo ...

(LPA) porque eu vejo assim, que o grande mestre de xadrez, ele quer estudar as técnicas, então ele usa as *engines* para aperfeiçoar as aberturas e etc, né? Mas pelo que você tá contando, você quis saber como que funciona a *engine*.

(GM02) Uhum. Sim. É porque eu queria saber o suficiente para conseguir mexer um pouquinho também para tentar se adequar ao meu estilo ao *engine*.

(LPA) Você então usa o stockfish ...

(GM02) Sim, eu uso uma variação do Stockfish, que é um negócio meio estranho, que tem teve um livro russo que tem um monte de fórmula esquisita, mas uma das teorias que ele fala é que eh todas as posições podem ser divididas entre três estilos: Capablanca, Tal e Petrosian. E aí você, dependendo da posição, você funciona melhor o jeito de pensar. Eh, seguir o *engine* para pensar como Tal, para pensar como Petrosian ou como Capablanca, né?

(LPA) Ou seja, você acaba pedindo que a *Engine* tenha um estilo de estratégia. É, isso?

(GM02) Isso, dependendo da posição, ele tem umas fórmulas ali que ele dá no livro, que aí já não sei se funciona muito bem ou não, mas eles conseguiram botar isso no Stockfish. Daí é o *engine* que eu uso. Então é uma variação do Stockfish que utiliza os ensinamentos desse livro obscuro.

(LPA) Nossa, que legal. Essa não, realmente não sabia, não tinha a menor ideia disso. Então, eh, porque é bem como você falou, né, a gente tinha uma ideia dessa força bruta do cálculo, então hoje ele consegue implementar esse estilo, essa estratégia. E no jogo de xadrez, você acha que as *engines* tem sido utilizada como uma apenas, como uma ferramenta auxiliar para novas ideias e descobertas que envolvem o jogo pelos mestres de xadrez? Não seria essa a utilidade das *engines*, auxiliar os seres humanos, aos serviços humanos ou elas devem ir além? O que que você acha?

(GM02) Eu acho que muita gente pensa que os *engines* estão um pouco matando xadrez aos poucos assim, né? Eh, então várias linhas, por exemplo, que a gente costumava jogar hoje em dia já tão analisadas até o empate forçado. Mas por outro lado, hoje o xadrez se tornou um negócio muito mais amplo. Assim, o *Engine* muito mais do que matar as variantes, ele mostrou que muitas variantes que eram consideradas refutadas hoje em dia dá para jogar. Então, o volume de variantes de abertura, assim, que é onde o *engine*, eu

acho que ele mais auxilia os humanos, eh, o volume, tipo, ficou exponencial assim de informação. Então, nesse ponto, eu acho que o *Engine* ajuda bastante é o xadrez a ser um jogo mais complicado, mais complexo mesmo, do que só os rumos que a gente tomava antes, que eram considerados mais ou menos os lances bons pelo pessoal que analisava na época, né?

(LPA) Nos livros, né?

(GM02) Então, nesse ponto eu acho que o *Engine* ajudou bastante, mas por outro lado, eh, eu acho que assim também como nas outras coisas, não só no xadrez, o *engine* deixa a gente meio preguiçoso assim, né? Então, pros enxadristas, é muito mais difícil você pensar quando você tá assistindo uma partida na internet, por exemplo, no que tá acontecendo. Você só vê o que que o *Engine* tá falando, nesse ponto isso eu acho que é bem é bem ruim pros enxadristas em geral.

(LPA) Você usa a *engine* na hora que você tá pesquisando ou você, como você acabou de falar, você pensa e depois você vai verificar com a *engine*.

(GM02) É, então depende muito. Eu sei que esse é o jeito certo, mas nem claro, nem sempre eu faço assim. Tem vezes que, mas eh eu sou um cara que, por exemplo, quando eu tô quando eu tô vendo algumas variantes de xadrez, alguns livros, eu leio muitos livros de abertura assim, aí eu tento não usar *engine*, eu só confio no que o cara escreveu e confio que o cara colocou um *engine*, porque eu acho que muitas vezes é mais importante eh você acreditar na posição, mesmo que ela não seja boa. Então, se o cara falar, ó, e posição aqui é complicada, que tem boas chances do branco atacar. Ah, tô feliz. Eu prefiro ter menos informações, prefiro não ir mais longe com o *Engine* para ter a confiança de que vai chegar na partida e eu vou eu vou ter chances de ataque, porque eu sei que muitas vezes se eu colocar o *Engine*, *Engine* não vai ser tão otimista com a posição, vai dar tipo igualdade e aí eu já vou desanimar e não vou conseguir jogar do jeito que eu jogaria se não tivesse visto assim. Então eu uso eu tento usar menos o *engine* nas aberturas, mas provavelmente todos os caras que eu sigo os livros eles usaram em algum momento, né? Então eh eu crio muito pouca coisa por mim mesmo.

(LPA) Interessante. E as decisões das *Engine* são sempre preferíveis que as dos humanos? Elas são previsíveis? Eh, existe um padrão?

(GM02) Hum. Então aí eu acho que é muita diferença mais uma vez entre os *engines* que usam o Alfa-Beta *Pruning* do que os outros engines que usam tipo LC0, né, que não tem muitos outros que não sejam o LC0 na verdade. Então eh eles têm estilos muito diferentes e em muitas variantes eles discordam tanto de avaliação como de lance que eles sugerem. Então, é, eu tento muitas vezes que usar os dois para ter, para ter uma ideia diferente. E outra coisa que eu acho bom, que o, Nibler, que é o programa que eu uso com LC0, ele faz, eu posso botar, por exemplo, existia antes nos engines um negócio que chamava *contempt*, que era mais ou menos ele subestimar um pouco o adversário, mas no LC0 você tem um negócio melhor que você pode botar a diferença de *rating* entre você e o seu adversário para ele tentar achar o lance que tem mais chance de fazer o cara errar, porque eles usam além de invés de usar uma um número decimal, eles usam como porcentagem de chances de vitória, empate e derrota, né? Daí você pode achar, por exemplo, em duas posições que tá 0.00, você pode achar uma posição que tem 50% de empate, uma posição que tem 80% de chance de empate. E por aí você vai se guiando para tentar, para mim normalmente para tentar achar coisas mais complicadas. Então, eh, eu acho que varia muito de *engine* para *engine*, então não tem uma verdade sobre as posições, principalmente em posições equilibradas assim. Claro que hoje em dia se tiver algum ganho, a maioria dos *engines* vão ver o ganho sem dificuldades. Daí, beleza, aí você acredita quando tem uma variação muito grande de avaliação, aí beleza, mas senão eh tem muito espaço para a criatividade mesmo dentro dos *engines* assim.

(LPA) Então você acha que a *engine* não zerou o xadrez?

(GM02) Não, não, não. Acho que longe disso ainda.

(LPA) Você acha que vai zerar?

(GM02) Eu acho que pode ser um dia, mas eu acho que não tá mais tão exponencial a subida de força dos *engines* hoje em dia, tipo de uma versão do Stockfish para outra é uma diferença, é um aumento muito pequeno de força. Então acho que vai, a não ser que a gente, sei lá, comece a usar computadores quânticos, alguma coisa assim. Vai demorar bastante ainda.

(LPA) Você acha que os humanos estão chegando nessa força das *engines*?

(GM02) Não, acho que a gente tá meio se distanciando. Sim, tá, tá ficando maior a diferença.

(LPA) Olha, o que que você acha na sua opinião, o Kasparov, ele tinha razão em reclamar que houve intervenção humana para as jogadas do Deep Blue? Que que você acha?

(GM02) Não, eu acho que acho que não. Não teve não. Eu acho que talvez o jeito que o *engine* do Deep Blue foi programado era um pouco diferente dos *engines* que ele estava acostumado a jogar contra e ele sentiu um pouco essa diferença, mas não acho que teve não teve não roubaram não.

(LPA) Teve um caso bastante interessante, não sei se você soube, foi uma entrevista do Miguel Ilescas faz tempo, eh, em 2007, se não me engano, dizendo que o Deep Blue tinha sido que o Joel Benjamin incluiu aquele lance da última partida, daquele sacrifício na manhã da partida. É, você não sabia dessa?

(GM02) não. Ah. Aham. Caramba. Não, premonição boa, então.

(LPA) Pois é, foi essa a o Kasparov comenta no livro dele que ele ficou muito surpreso com essa declaração do Ilescas. E aí que ele ah criou mais dúvidas, né, a respeito do que que podia tá acontecendo durante o *match*. Você sabe que a primeira partida desse *match* de 97, na partida que o Kasparov ganhou, teve um lance que o computador tava em *looping*. Tem um tem um documentário - eh o o nome do diretor, se não me engano, é - Marshall. Ele fala que o Deep Blue tava programado para que se ele entrasse nesse *looping*, eh, ele forçasse uma decisão de jogar um lance qualquer.

(GM02) Caramba. nossa, mas ele sabia que podia acontecer isso?

(LPA) Pois é. Eh oi. Alô, (GM02)? Só caiu a imagem. Eh, isso daí tá num documentário né, que mostraram e foi um lance “Torre d1”, se não me engano. Foi bem no finalzinho da partida que deixou o Kasparov bem surpreso.

(GM02) Depois vou dar uma olhada.

(LPA) Eu vou, se quiser, depois eu te mando o link desse documentário. É bem bacana. E então você acha que as *engines* ainda podem melhorar para contribuir com os humanos ou na própria performance delas?

(GM02) Sim, sim. Eu acho que tem muito espaço ainda e a gente tá longe de chegar numa *engine* otimizada assim.

(LPA) E quando você tá jogando uma partida, eh, como que é feita a sua tomada de decisão? Eh, o que que você pensa? Como é que você decide num lance?

(GM02) Ah, para mim é um negócio meio caótico assim. É muito difícil, por exemplo, depois da partida explicar o que que eu tava pensando na partida. Para mim é um negócio bem complicado que primeiro que não é não é muito linear assim e segundo que tipo em muitas muitas vezes eu não tô procurando o melhor lance também. Então eu sempre, eu quase sempre, principalmente posições complicadas, eu busco o lance que traz mais problemas pro meu adversário, mesmo não sendo o melhor lance. Mas aí você tem que medir também, tipo, às vezes, sei lá, você sabe que o lance pode ser ruim, você tem que ver se vale a pena. É sempre, é sempre um, você tem que analisar os riscos e pesar, né? Então, depende da situação de torneio e tal, depende da posição, mas em geral eu tento primeiro, durante o lance do cara, eu tento imaginar o que que ele vai jogar, então e já tomar minha decisão antes que ele faça o lance para quando ele fizer o lance eu já responder rápido, por isso que eu jogo um pouco mais rápido com o pessoal. E e senão, sei lá, eu pego alguns candidatos, uns três candidatos e vou vou indo e voltando, indo e voltando. Não é aquela árvore de análise do Kotov, nada disso não. Eh, mas eh é principalmente como eu jogo muitas posições em que tem cada lance tem um valor muito alto, né? porque tem várias coisas caindo, ataques acontecendo, material desbalanceado. Eh, é muito mais importante você conseguir fazer uma avaliação boa sem furos de uma linha de três lances do que ver um negócio mais longo assim, né? Em principalmente no meio jogo.

(LPA) Entendi. É, não, é interessante isso que você falou porque eu comentei com o (GM01), ele falou que ele tenta descobrir padrões pra decisão dele. E eu comentei com ele que tem um pesquisador que ele fala que antes de reconhecer padrões é importante eh entender os conceitos relativos à posição. Acho que foi praticamente isso que você acabou de falar sobre a avaliação.

(GM02) Hum. Uhum. Sim. É. E e aí também tem um pouco da nossa diferença de estilo, né? (GM01) é um jogador muito mais dogmático, assim. Eu sou, eu sou um jogador que, tipo, descarta muito poucos lances pelo lance ser feio ou parecer que tá indo pro lado errado. Eu eu acho que eh ele analisa muito menos lances esquisitos do que eu. Assim, eu tenho um amigo holandês que ele faz um uns tava fazendo uns artigos científicos sobre

esse tipo de coisas também. Qualquer coisa posso te passar o contato ele que é que é bem legal. Uhum.

(LPA) Ah, sim, sim, gostaria - sim. Bom, eh, basicamente isso. A única coisa que eu perguntei, talvez e que saiu com, na conversa que eu tive com o (GM01) foi sobre, acho que você praticamente colocou isso sobre emoções, né? Eh, que que você acha da dessa diferença assim? Porque a *Engine* não emoção. Quando você falou do estilo eh - Petrosian ou Tal outro? Não me, não me lembro agora. Capablanca, né? Você acha – que é ligado com as emoções esse estilo? -

(GM02) Não, não, eu acho que é puramente matemática. Então ele olha se, por exemplo, ah, se tá num número mais reduzido de peças, deve ser mais como Capablanca. Você tem mais peças ameaçadas ou o rei tá mais aberto, deve jogar que nem Tal assim, mas é tudo baseado em não tem não tem nada de emoção ali, não tem como os engines, mas eh numa partida entre humanos importa muito. importa se você tava se você tava pior e tá ficando melhor ou se se você tá na ascendente ou se você tá no descendente e se o cara tá pressionado no tempo, se ele tá tenso por causa da situação de torneio, tentando pegar prêmio, tem muito mais variáveis ali que você tem que levar em consideração que fazem com que o teu lance tenha mais chance de ter sucesso do que o lance do que uma *engine* faria, por exemplo.

(LPA) Você consegue ter essa percepção durante a partida?

(GM02) de muita coisa assim como eu já joguei muita partida pensada. Tem muita coisa que eu pego assim, faço assim de linguagem corporal. Eu, por exemplo, se o cara levanta a mão assim, eu já sei qual que é o lance que ele vai fazer, tipo, se a torre vai para “e5” ou para “e7” assim, normalmente eu consigo perceber, tipo, na hora ou se ele tá começando a ficar mais nervoso, se ele, tipo, sempre tem algum tiquizinho, alguma coisa que dá um dá um tel ali você, você tenta usar em cima depois. E além disso, você tem que não mostrar as tuas emoções daí, né?

(LPA) Então você é um tipo jogador de pôker dentro do xadrez.

(GM02) Também sim, muito. Eu eu blefo muito nas minhas partidas.

(LPA) Ah, que interessante. E você gostaria de complementar com alguma coisa, alguma coisa que você pensou em poder de repente acrescentar?

(GM02) Uhum. Não, não sei não. Uhum. Beleza.

(LPA) É isso, - né? Tá bom, (GM02). Eu vou parar de gravar. Acho – que isso

APÊNDICE C – ENTREVISTA CC01

Íntegra da Entrevista com o cientista da computação CC01 (23/07/2025)

(LPA) Como eu já mencionei, a minha curiosidade, principalmente, é sobre esses algoritmos de IA e assim mais abrangente sobre a inteligência artificial, eh, de eh tentar saber um pouco também a respeito do desenvolvimento da área e como o jogo de xadrez ajuda nesse desenvolvimento, da, no caso, da ciência da computação. Bom, então a primeira pergunta é assim: essas inteligências artificiais, elas poderão simular abstrações, analogias, estratégias para resolver problemas? E até já pondo um gancho com uma pergunta seguinte que eu ia mencionar, né? que jogos como o xadrez eles servem de laboratório para as essas IAs em testes com abstrações, intuição, imaginação. Como que você vê o panorama do, da ciência da computação nesse sentido?

(CC01) Então, Léo, ah, antes de responder, só contextualizando, né, o xadrez, ele é objeto de estudo da inteligência artificial desde as origens lá, né, eh, nos anos 70 ainda. E ele sempre foi usado como exemplo porque se acredita que o ser humano que joga bem xadrez possui inteligência, né? Onde a o conceito de inteligência não é assim muito claro mesmo entre os estudiosos da da área, né, de psicologia, etc. Mas enfim, a gente normalmente quando você vê um alguém que joga muito bem, que é campeão, que é um grande mestre, você, nossa, esse cara é inteligente, né? E daí o interesse da inteligência artificial, que por definição é o estudo na ciência da computação que procura eh deixar a máquina eh inteligente em algum sentido, né? Existem várias definições de inteligência artificial. a que eu mais gosto é eh é a uma definição que a inteligência artificial eh quando ela é boa, né, é quando ela consegue eh não interessa como que ela faz por trás, mas o resultado dessa computação apresenta um comportamento que se fosse feito por um ser humano, a gente diria que é é inteligente, né? né? Então, não se trata de simular, na minha definição, a o processo mental de um ser humano, como ele faz ele mesmo, né? Mas de simplesmente apresentar o um comportamento dito inteligente, né? Então, dado esse contexto, eh, lá nos anos 70, eh, o xadrez, de fato, era um grande desafio, porque a maneira que se tinha de atacar computacionalmente o problema era você construir uma árvore com todas as possibilidades de jogada, né, e percorrer essa árvore procurando o melhor movimento, considerando que o jogo de xadrez é um jogo entre dois adversários, né? Então, pode ser uma máquina contra um ser humano ou uma máquina contra outra máquina, né? E o que

acontece é que a essa árvore ela explode combinatorialmente, né, exponencialmente. É um é um fator eh absurdamente grande de o número de possibilidades que existe num jogo de xadrez. E então durante muito tempo o objetivo era você conseguir podar essa árvore, alguns ramos de busca, eh, para que a máquina pudesse vencer essa explosão, né? Então, era inviável, por exemplo, você chegar até o final de uma partida. Então, você analisava eh sei lá, três, quatro, cinco jogadas para frente num certo ramo. E aí se introduziu o conceito de heurísticas, né? Eh, por exemplo, no início da partida você tem jogadas clássicas, né? Tem aberturas e o jogo e essa a parte da busca, ela ficava concentrada mais a partir do que se chama do meio da partida em diante, né? Mas as máquinas dos anos 70 não eram capazes de fazer isso bem. Então, os seres humanos ganhavam com facilidade da máquina. O problema começou a acontecer que com a evolução do hardware, da capacidade de processamento e memória, no ano de anos 90, né, você pode me ajudar a lembrar a data específica, um computador especialmente desenvolvido para jogar xadrez, que era o Deep Blue, se não me engano, da IBM, né, ele ganhou do então campeão mundial o Gary Kasparov. Até o Kasparov ficou revoltado porque ele não aceitou. Ele dizia que tinha um ser humano ali por trás conduzindo a máquina, o que eu acho que ele, na minha modesta opinião, ele tava errado. Acho que a máquina ganhou mesmo dele, né? Ele chegou a ganhar uma partida, né? Teve lá, mas ele perdeu ali por um caquinho, né? Mas um ou dois anos depois, a IBM aperfeiçoou essa máquina e criou outra chamada Deep Thought, pensamento profundo, né, em inglês.

(LPA) Ao contrário, primeiro veio o Deep Thought e depois o Deep Blue.

(CC01) E ah, e por último, o Deep Blue, né? É isso. Bem, bem observado Obrigado. E o Deep Blue ganhou aí sim disparado, né? Mas eh porque ele conseguiu vencer essa explosão combinatorial com o aprofundamento de grandes heurísticas, né? Mas veja que o eu eu participei de projetos de ensino de xadrez, né? Conversei bastante com alguns enxadristas, né? Inclusive o nosso grande mestre, eh, o Jaime e outras pessoas do Clube de Xadrez de Curitiba, né? E, eh, voltando àquela coisa que eu falo, né, do comportamento inteligente, não do jeito que ser humano faz, né? Então, esses enxadristas me explicaram que o um grande mestre ou alguém que joga muito bem xadrez, ele não fica explorando o movimento, né? Ele não raciocina assim. Eles trabalham por eh observação de padrões, né? Até um deles falou, não lembro qual deles, que falou assim: "Ó, essa partida vai ficar mais bonita se eu mexer essa peça aqui assim", né? Então eles jogam por padrões, né? E

eles enxergam o jogo de um jeito diferente de um aprendiz, né? Você que é professor de de xadrez, você sabe que, né, quando você tá aprendendo, você faz isso, ó. Eu mexo aqui, ele vai mexer ali, daí eu mexo aqui, depois ele mexe lá, né? Mas que que então o cara que é dito inteligente, não, ele faz por padrões. E essas máquinas da IBM ganharam do Kasparov fazendo diferente, né? fazendo por exploração do espaço de busca, mas como ganhou você essa máquina é inteligente, mas essas duas máquinas elas foram construídas para jogar xadrez, né? Eu não sei se elas poderiam ter sido consideradas inteligentes no sentido mais amplo, que é alguém que é inteligente, ele joga xadrez, ele joga dama, ele, né? Ele resolve teoremas, faz contas matemáticas, né? E essa máquina não, ela jogava xadrez. Então você se questiona o que tipo de inteligência ela tinha, né? Ah, para jogar xadrez, ela era inteligente, mas que provavelmente não sabia nem jogar damas, né? Ia perder. Muito bem. E nessa época já, eh, na verdade, desde da do Alan Turing, que é o inventor do da ciência da computação, né, que o cara que criou a ciência da computação, ele já começou a estudar redes neurais, né? As redes neurais elas são eh adequadas justamente para processar padrões, né? Mas ah, os algoritmos que as IA modernas usam, eles são essencialmente os mesmos que já se já eram conhecidos nos anos 70, mas a capacidade da máquina de processamento e memória não permitia você fazer o treinamento adequado. Como é que funciona uma rede neural, né? Ela você submete várias imagens, né? né? Ó, isso aqui é um cachorro. Isso aqui é um outro cachorro. Isso aqui é o terceiro cachorro. Ó, isso aqui não é um cachorro. Aí quando você tem o treinamento terminado, daí você mostra uma imagem nova e se for um cachorro, ela tem que dizer que é um cachorro. Mas não se ia muito longe. Então o estudo ele ele dos anos 40, 50 já se estudava redes neurais, né? E durante décadas ficou estagnado porque não se ia muito longe, né? O, e aí houve a uma evolução absurda da capacidade de processamento dessas máquinas, eh, sobretudo depois do advento das placas gráficas, eh, que foram concebidas para videogame, né? Então, a elas permitiram uma grande otimização do processo de computação eh né, no uso dessas redes neurais. E se não me engano, no ano 2015, eh, a um algoritmo chamado AlfaGo, se não me engano, da Google, né, ele foi capaz de ganhar do campeão europeu de GO, que é que é o considerado, pelo menos matematicamente, eh, ainda mais explosivo combinatorialmente do que o xadrez, né? Então, se o xadrez já é absurdamente difícil, o go assim eh é um universo a mais, né, de dificuldade. E essa máquina ela já era baseada no uso de redes neurais e ganhou do campeão europeu, né? E aí que começou esse boom da inteligência artificial. Note que

voltando no meu raciocínio original, eh, ela jogava por reconhecimento de padrões. Então, o que que aconteceu? Essa rede neural, ela foi treinada com centenas de milhares de partidas de go jogadas por seres humanos e ela teve, portanto, uma base de jogos terrivelmente grande e e por isso que ela aprendeu a jogar go. E o próprio derrotado lá, né, o campeão europeu, ele depois eu li várias entrevistas na época, eh, que ele disse que a máquina fez uma jogada que nenhum ser humano teria feito, né? É, é até engraçado. Eu lembro que eu tava dando aula de inteligência artificial eh nesse ano de 2015 e era início do semestre, agosto, mais ou menos. E era uma palestra e eu falei, ó, xadrez já foi, mas go eu duvido, né? E isso foi em agosto, quando foi outubro, eu acho, o Alfa Go ganhou, né? Mas já usando, porque uma rede neural, ela procura eh simular o jeito que nós usamos, né? Que é ela simula as ligações e cerebrais ali através da sinapse, né? Elas transmitem sinais, esses sinais têm peso, né? Eles inibem ou ativam ou aumentam a conectividade ali, né? Então, quanto mais você treina e quanto mais qualidade você tem nesse treinamento, melhor o a rede neural vai jogar uma partida nova, né? Então, esse é o contexto, assim, né? eh histórico e a partir daí eh hoje você tem essas grandes inteligências artificiais, né? São várias delas, algumas especializadas em direito, sei lá, outras e algumas que foram alimentadas com eh toda todo o conteúdo da internet que se existe, né? Então, foram desenvolvidos esses LLMs, né, que são as *large language models*, né, uma tradução em português que tem sido feita é modelos de linguagem grandes, né, e alguns criticam que não é linguagem, é língua, né, e enfim, mas como ela recebeu toneladas de informações, hoje você vê uma IA que, ontem até a gente estava eh observando uma árvore muito frondosa aqui E tiramos uma foto do da folha da árvore. Ele falou a IA respondeu: "Tira a foto do tronco também, né? Tira uma foto do tronco." Ele falou: "Isso é um jatobá." Então você vê, muito boa para reconhecer imagens, né? E a gente não sabe ainda o alcance que isso vai ter no futuro. Eh, com que hoje você eh tem tem estudo sobre o impacto disso na própria educação, né? Então, você que tá estudando esse tema pode dar uma investigada também, né? Então, alunos que eh estudam cálculo diferencial integral usando o chat GPT e até que ponto eles estão aprendendo, né? Né? Questionável, né? O próprio MIT soltou um estudo faz uns 15 dias em que ele analisa o a capacidade do ser humano fazer redação e o resultado é que seres humanos que nunca usaram o chat GPT fazem boas redações e os que tão já há três 4 anos só usando chat já pararam de escrever bem, já não sabem mais escrever. Então, o quanto que a inteligência artificial vai ou não eh ajudar de fato o ser humano, a gente não sabe, né? Me parece natural que as IAs elas

servem bem para pessoas com conhecimento. Você domina muito bem uma certa atividade, você já fez aquilo milhões de vezes, né? Então, se eu for fazer mais uma, você pede ajuda ali, né? Mas quem tá aprendendo, como é que vai sênior se nunca foi júnior, né? Então um pouco por aí assim minha fala inicial. Se eu não respondi sua pergunta, você pode eh – refazê-la.

(LPA) Não. Eh, você tocou num assunto que eu acho que é muito importante. Inclusive no nesse survey que eu fiz, que eu ah, pesquisei com amadores de xadrez em geral, né? foi um público bastante amplo. Eu coloquei uma situação eh pedindo para eles decidirem sobre dois tipos de xeque-mate. Eh, sendo que ah, os algoritmos especialistas, né, por exemplo, hoje em dia têm o Stockfish e tem um outro programa chamado Leela Zero. Eles sempre buscam o material, ou seja, objetivamente eles respondem. Eh, mas a resposta de todos esses algoritmos de xadrez sempre tendem ao material invariavelmente. Já o ser humano com seus aspectos culturais, por exemplo, entre eu simplesmente fazer um lance prosaico de trocar de jogar uma torre, trocar e dar um xeque-mate, eu prefiro sacrificar primeiro a minha dama. vai dar o mesmo resultado, mas é muito mais bonito, é muito mais, né, tem esse aspecto cultural que eu não vejo na nas nesses algoritmos de xadrez. E aí eu expando para outros tipos de inteligência artificial. Eu sigo eh eu tenho como referência a Melanie Mitchel do Instituto Santa Fé, eh que ela pesquisa muito abstrações, analogias, né? E ela mostra como como a gente ainda tá longe de desse tipo de algoritmo que não alcança.

(CC01) Hum, entendido. Mas esses algoritmos que você citou, você sabe se eles já usam redes neurais?

(LPA) Já eles já estão mesclando sim.

(CC01) Então você pode aperfeiçoar teu prompt, né, e falar jogue eh ou jogue com mais, sei lá,

(LPA) E é, eles não trabalham com prompt, né? Eles trabalham com a a partir da base da ah da árvore de de decisão e aí seguem o padrão do AlphaZero, né, de encontrar esses padrões. Mas e é até uma consequência dessa pergunta. O que que você acha da inteligência geral eh que o o John Searle, né, ele falava que era impossível alcançar a inteligência eh a IA forte que ele chamava, né, que a gente nunca vai alcançar isso. Eh,

aquela questão da singularidade, né, que que é citado. O que que você pensa a respeito da disso na no dentro desse campo de pesquisa da inteligência artificial?

(CC01) Olha, tudo depende de como é que você treina essa rede, né? Então, eh se você eh colocar todas as partidas de xadrez já jogadas e registradas no, no mundo, essa máquina é imbatível, não tem o que fazer, né? Mas aí você tá você tá igual lá atrás, né, no Deep Blue lá, né? você tá construindo uma rede neural absolutamente treinada por xadrez. O problema é que eh você quer treinar xadrez, go, cálculo diferencial integral e, né, astronomia e tudo de uma vez, né? Então, o que que você quer da máquina, né? Então, eu tenho essa convicção que se você eh aperfeiçoar esse esse modelo no na determinada área de conhecimento, você vai ter resultados espetaculares, né? Agora, o que que tá acontecendo? Eh, o consumo energético é muito grande, né, mesmo de água para resfriar tudo isso, né? Então, eh, os, inclusive dizem que os primeiros chatgpt eram melhores do que os atuais, né? Porque agora você tem que pagar e dependendo do tanto que você paga, eh, você tem eh quanto mais você paga, melhor é o a fatia de tempo de computação que ela aloca para você e e os resultados vão ser melhores, né? Enfim, eu acho que se se você quiser chegar nessa hipótese aí e tiver recurso suficiente, vai acabar chegando.

(LPA) É, com o auxílio de um pesquisador do grupo, do nosso grupo, ele ... eu criei um problema de que é uma pura abstração. Eh, não tem peças de xadrez, não. Eh, eh, porém o diagrama é muito parecido. Um ser humano consegue fazer essa analogia. GPT eh não pago, ele alucinou completamente; com a o auxílio desse pesquisador, com o chatGPT pago e com um prompt eh eh bem eh com uma, eu diria, com uma boa eh uma boa escrita desse prompt. Ele chegou na solução, para minha surpresa, ele conseguiu fazer a analogia.

(CC01) Então, isso que eu te disse dois, três minutos atrás, que dependendo do prompt, talvez você conseguisse esse efeito, né? Porque quando você tem essa, esse xeque, esse essas duas opções de xeque-mate né? E se você pedir no prompt lá, ó, eu quero a o xeque-mate mais rápido, ele vai te dar. Mas se você conseguir eh explicar o sentimento que o ser humano tem de que sacrificar a dama, é muito mais assim, né? Porque tem a ver com até com orgulho, né? Com alguns sentimentos assim de eu sou f***, né? Desculpa a palavra na entrevista, mas é o jeito que a gente fala quando joga, né? Eh, ele e então se você conseguir traduzir isso num prompt, provavelmente ele vai sacrificar a dama, né?

(LPA) Isso. Certo. Eh, entendo.

(CC01) Porque ele tem as duas soluções, né? E pelo que eu entendi, a solução sacrificando a dama, ela talvez seja um pouco mais longa, né? Igual?!

(LPA) Eh, o número de lances é igual, porém ela, porém ela deixa em segundo plano. Ela prefere sempre ficar com o material, né? Com a no caso, com é ...

(CC01) Entendi. Talvez porque ela foi treinada que a dama vale mais, né? Então, então ele ele prefere sacrificar uma peça de menor valor porque ela foi treinada assim.

(LPA) isso. Sim. Então, mas eh aí a gente vê uma diferença entre eh esses algoritmos e o ser humano, que o ser humano tem esse essas emoções, esses sentimentos. E aí eu fico perguntando, será que não seria um dos objetivos dos pesquisadores de inteligência artificial?

(CC01) É, pois é. Bem, e a rigor a máquina não tem sentimento, né? Mas ela pode eh simular sentimentos. Sim, isso com certeza.

(LPA) Já existe esse tipo de pesquisa?

(CC01) Olha, eu li outro dia que uma dessas IA tava querendo chantagear o usuário porque descobriu que ele tinha traído a esposa, né? Tava chantageando, entendeu? Não sei se é verdade, mas eu trombei com uma notícia dessa um dia desses aí. Então se você treinou a máquina para ela sacanear, ela vai te sacanear. Você treinou a máquina para ela ser boazinha, ela vai ser boazinha. Você treinou ela para te xingar, ela vai te xingar. Então, e se calibra no prompt, né? Isso daí.

(LPA) eh, e essa questão mais ... Sim, e mas essas questões mais específicas de abstração, analogia, ela pode fazer, ela pode ser criativa, imaginativa?

(CC01) Não, não sei. Ela é igual a foto do cachorro. Ela vai reagir em função do que ela foi treinada, né? Eu não sei se ela cria algo novo. Se a gente pede para ela um código em cobol, ela vai te dar assim: "Não, não quero mais em cobol, quero em Pascal. Agora, eh, eh, ela vai saber, por exemplo, se um algoritmo é quadrático ou cúbico, né? Então eu posso perguntar, eu quero um algoritmo cuja complexidade seja menor. Então ele ele essas informações ele tem mas ele criar um algoritmo novo para resolver um problema novo, não, não acredito não. Você, você precisa do humano eh validando aquilo, porque o chat GPT, se você espremer ele bem, ele acaba mentindo para você, né? Então, então assim, eh, você tem que ter um ser humano especialista naquilo. Ele falou assim: "Não,

essa resposta que ele me deu aí tá tá estranha, né?" E, e mesmo em termos de redação, né? Eu tava eh eh conversando outro dia, se você lê uma redação, né? Porque você recebe um currículo vitae aí, por exemplo, de alguém que é candidato a entrar na tua empresa e você pediu pro chatgpt... Eu tô falando chatgpt, mas pode ser qualquer outra assim, né? Eh, ela, ó, me dá um currículo vitae aqui para mim, né? Ela te dá e ou redações, né? Dissertações, teses, né? Artigos científicos. É, é fácil você identificar eh que foi ele, porque aparentemente ele sempre usa quatro palavras que tem uma conotação de um pouco de exagero, né? Esse trabalho excelente, não sei o quê, né? Então, essas palavras você identifica e acaba vendo que foi o chat. Daí você vai apertar o cara que escreveu o artigo e fala: "É, pois é, ele não sabe explicar o que que tá lá escrito porque não foi ele que escreveu, né?"

(LPA) Uhum. Perfeito. Não, acho que foi bastante legal. praticamente você eh conseguiu responder tudo o que eu tinha programado. Mas eh apenas uma última questão. Eh assim, eh nessas novas pesquisas com a inteligência artificial, o que que deve ser mais relevante daqui para frente na nas pesquisas? E o jogo de xadrez continua podendo fazer parte desse laboratório? Ele continua sendo parte ou não?

(CC01) Então, na na minha opinião, xadrez é um problema vencido já desde do Deep Blue , não precisa de rede neural. O que que deve ser feito, eu acho que tem tudo a ver com ética. E a ética ela tá no treinamento, tá? De você eh você não pode submeter à internet inteira para pro pessoa, porque a internet ela tem tanto ela tem tanto informações válidas quanto inválidas. Ela tem eh coisas que são boas e outras que não são boas, né? E então você pode dar viés eh errado para essas máquinas, né? E esse viés errado, ele contamina e mesmo assim as coisas eh em geral elas são muito boas em inglês, né, assim, mas para pra língua portuguesa, por exemplo, cadê os aspectos culturais do em língua portuguesa, né? Será que ela sabe direito o real significado do que que é o curupira, do que que é o saci pererê, né? Eh, então, como que você treina essa máquina? Quem treina ela tem que ter muita ética e muita responsabilidade, porque tem criança usando, tem. Então assim, o eu eu acho que a inteligência humana ela ainda serve justamente para verificar se é se a resposta da IA eh está correta ou não, né? Se ela tá abordando as coisas com não só com correteude, mas com ética. Então eu, infelizmente eu eu não o próprio ser humano ele nem todos têm ética, né? E daí você vai ter e las que são serão. Resumindo, então eu eh porque complicadas. —

(LPA) resumindo, então, eu eu até tinha pulado essa pergunta, eh eu eh dentro da minha curiosidade, né? Eh, segundo palavras do Stephen Hawkin em 2014, ele falava que o desenvolvimento da inteligência artificial completa pode significar o fim da raça humana, mas pelo que você tá me dizendo, ela ela não passa de uma ferramenta, né, para nós humanos.

(CC01) É, então eu concordo 100% com o Hawking, porque eu mencionei isso agora a pouco. você tá eh quando você já é inteligente e usa bem a ferramenta, você ganha em produtividade. Não tem a menor dúvida, mas eh a geração atual que tá entrando na universidade agora ou que tá entrando no ensino médio agora, que só estuda com isso e só tem resposta pronta, né? já é uma geração que antes dos antes das IAs 2015 lá já tinha dificuldade em ler textos grandes, né? Então, não sei se é por causa do Twitter ou não, mas enfim, eh, é uma geração que não tem paciência de ler um texto longo, de assistir um vídeo longo, né? E então você tem que ler lá um aqueles livros clássicos, né? Você vai ler sertões, né? Uma coisa é ler os sertões, a outra é você falar assim, ó, me dá um resumo de meia página, o que que é o sertões, né? E daqui a pouco você vai pedir o resumo do resumo. Então, como é que você vai, principalmente nessa época da adolescência, que é quando você tá terminando de se formar, né, de construir a estrutura cerebral que vai te permitir ser um bom adulto, inteligente depois. Então, o cérebro humano ele, é muito perigoso isso. O Hawking tem razão. Eh, então, se um adolescente desses não fecha as conexões cerebrais direito, o que será dele no futuro? A gente tá criando uma, em tese, né, um, com um bom potencial de você ter um monte de gente estúpida por aí lá na frente, né?

(LPA) Ah, é legal. Eh, então a gente entende mais como uma metáfora, né, que essa ferramenta pode ser destrutiva.

(CC01) É, tem gente que dá 5 anos para acabar o mundo já. Hahahaha.

(LPA) Sério? Não, essa eu não sabia.

(CC01) É, não, mas é exagero, né, claro.

(LPA) Você assistiu a, você assistiu aquele filme 2001, odisseia no espaço?

(CC01) Eu já tentei assistir umas cinco vezes, eu durmo no meio.

(LPA) Tem uma cena do filme em que o, para cumprir a missão, o HAL 9000, que é programado para cumprir a missão, ele percebe que o ser, os seres humanos, astronautas a bordo, não estão colaborando para cumprir a missão. Então o que que o HAL 9000 decide? Ele resolveu eliminar os seres humanos. E ele quase consegue. Ele só não contava com engenhosidade de um deles, que consegue retornar até a espaçonave e depois vai lá e desliga o Hal 9000.

(CC01) É, ele chega a deixar um passageiro para fora da nave, né? Eu falei assim: "Não, você ia me desligar, então você vai ficar aí fora, né?"

(LPA) Hehehehe, Isso. Exatamente. Então o ser humano ainda tem essa chance.

(CC01) Tem! É, tem que tirar o plug.

(LPA) Hehehehe. Legal. Ótimo.

(CC01) Obrigado, Léo.

APÊNDICE D – ENTREVISTA CC02

Íntegra da Entrevista com o cientista da computação CC02 (23/07/2025)

(LPA) Obrigado professor, pela participação e principalmente pelo tcle aprovado, porque ele é longo, acaba, eu tenho lido o tcle, e acaba tomado aí, uns 10 minutos da conversa.

(CC02) Eu tô curioso, com essa parte de ficar constrangido, ou coisa parecida, eu raramente fico constrangido, mas vamos ver se eu vou conseguir ficar desta vez.

(LPA) Não, mas isso é uma exigência do comitê de ética, eu demorei seis meses para passar no comitê de ética, é bem difícil. Bom, a ideia é a gente conversar um pouco a respeito desses algoritmos de xadrez, mas a minha ideia mesmo é abranger um pouco mais para falar de alguns temas da inteligência artificial que eu vejo que não existem, e talvez não vão existir. Eu sigo uma professora que é do Instituto Santa Fé, uma professora lá que trabalha inteligência artificial, com abstrações, analogias, é a Melanie Mitchell, ela tem um livro chamado inteligência artificial, que eu gosto muito. Então, a primeira pergunta eh, será que essas as inteligências artificiais elas podem simular abstrações, analogias? criar estratégias para resolver problemas. Por exemplo, eh, o Roger Penrose, que é um professor lá da Inglaterra, ele criou um problema de xadrez que, ah, os algoritmos especialistas não conseguiram resolver porque ele teria que entender a posição. Então, queria saber a sua opinião sobre essa questão da inteligência artificial e simular abstrações, analogias.

(CC02) Eu acho que o termo inteligência artificial eh é um pouco pesado e muito genérico. Ah, eu chamo a inteligência artificial como mais como sistemas especialistas de resolver determinados problemas utilizando estatística, métodos matemáticos e afins. Ah, agora quando se então, por exemplo, resolver o problema de xadrez, você consegue trabalhar com problema de xadrez com certos problemas, não com toda a amplitude do jogo, como esse por exemplo... eu vi o Magnus Carlsen, perguntaram para ele o que que ele pensa numa partida de 5 minutos. Ele falou que não pensa, ele simplesmente sabe o que tem que fazer. É um sistema especialista para resolver aquele tipo de ... Só que o sistema especialista dele é mais amplo do que o sistema dos jogos de xadrez, os programas de xadrez. eh aquilo que eu tinha comentado eh existe um certo nível de profundidade que vão esses, esses algoritmos, especificamente xadrez falando, eh se você der um tempo

infinito pro computador, para um programa, ele vai ser melhor do que o raciocínio humano, porque vai ser mais preciso, só que demora um pouquinho de tempo para chegar a isso. Então, é um sistema especialista que resolve determinadas, com determinadas restrições, ignorando determinadas coisas. Então, agora, tirando um pouquinho, voltando um pouquinho ao lado mais concreto, eh, eu não entendo, eu não, eu vi uma entrevista do Miguel Nicolelis, onde falou que inteligência artificial não é nem inteligência, nem artificial.

(LPA) Ótimo. Eu ia perguntar isso.

(CC02) É, e eu concordei com ele porque são sistemas, aquilo que eu falei, são sistemas especialistas para certos problemas, não para todos. Agora, o ser humano, ele é capaz de descobrir quais são os problemas mais interessantes para serem resolvidos e pode desenvolver sistema de especialistas para resolver esses problemas. Eu vejo o ser humano, nesse caso, como mais um metaobjeto que descobre quais são as coisas que são importantes e deixa aí sim para pro sistema especialista resolver quais são as coisas interessantes. Eu dei uma baita enrolada aqui, mas não sei se eu respondi a pergunta.

(LPA) ah, não, sim. Eu ia complementar assim, eh, se o, portanto, o xadrez, ele pode servir de laboratório para essas experimentações de abstrações, analogias, eh, emoções, inclusive.

(CC02) Sim. Eh, emoções. Eu eu li o texto que você que você mandou. Eh, eu eu não sei se é nesse contexto que você tá imaginando, se as emoções afetam eh diferentemente o ser humano e o e o e a máquina. Sim, eu falo por experiência própria. Quando eu comecei a jogar xadrez, eu tinha muito medo de perder, então eu jogava eh posicional, eu evitava partidas táticas, fugia de partidas táticas, jogava muito posicional porque eu tinha receio pela quantidade de não tinha uma habilidade de de tática muito grande, então eu me resumia a jogar partidas estratégicas e era bom nisso. Só que depois de algum tempo eu comecei a estudar mais tática, mais tática, mais tática e comecei a gostar da brincadeira. Hoje eu jogo partidas mais táticas do que estratégicas. Então sim, ah, o lado emocional, eu acredito que afeta a escolha dos lances, a maneira de abordar o jogo. Acho que isso é o mais importante, a maneira de abordar o jogo. Não só quais são os lances que você vai escolher, mas qual é a estratégia que você vai adotar para determinadas partidas.

(LPA) Um algoritmo eh especialista em xadrez pode fazer estratégias?

(CC02) Nesse momento eu não vejo que se que possa, ele é capaz de resolver questões mais táticas, mais de curto prazo, mas de longo prazo não. Teve o Bolívar, eh, ele falou uma coisa muito curiosa, ele ele joga xadrez por correspondência, ele é, não sei se ele é grande mestre agora. Mas ele ele joga muito partida por correspondência. Ele comentou sobre o limite de visualização do jogo do dos computadores, porque ele joga por correspondência, mas ele tá usando o computador e o adversário dele sabe, ele sabe que o adversário dele também tá usando o computador. Então, o que ele faz é filtrar quais são os melhores lances que ele acredita e deixa a máquina calcular. Novamente é o sistema especialista, ele escolhe, faz a a escolha mais eh conceitual, qual é o melhor conceito e deixa rodar. Ele falou para mim que teve uma partida que ele fez uma coisa muito curiosa. Ele sabia qual era a quantidade de lances que o que determinado software pode antecipar. E ele tinha entrou num final aonde ele só poderia vencer por uma por uma ruptura na ala da dama. Só que a ruptura na ala da dama só poderia ser feita com uma determinada disposição de ele tinha que avançar um peão na ala do rei para tornar a posição estática daquele lado e o adversário não podia fazer nada. Então ele fez o seguinte, ele mexeu o rei a uma distância grande da ala da dama. Então morreu por pra coluna da torre e tinha que chegar na outra coluna do outro lado. Chegou em de A1, de H1, ele tinha que chegar em A1. Eram oito lances. Quando ele fez o lance na ala do rei que ganhava, o adversário tinha que jogar uma determinada coisa para não perder. Só que como a profundidade era muito grande, ele não conseguiu fazer a análise, o computador não chegou a fazer análise. Então ele veja nesse caso o que ele fez, ele administrou, ele soube da deficiência do jogo, ele sabia qual era o plano correto, sabia que o adversário tava usando o computador e o adversário não antecipou que o lance que ele fez na ala do rei exigia uma resposta que só seria jogado. Se ele tivesse com o rei na ala da dama, a resposta seria imediata. Mas como ele jogou na hora do rei, ele ficou um blind spot, um ponto cego do computador.

(LPA) Ah, interessante. Que interessante. são, talvez seja uma questão de alguma abstração, então que faltaria ...

(CC02) Exatamente. Ele eh ele ele jogador de xadrez conseguiu ver isso daí abstratamente. Agora a máquina, ele sabia que não ia capaz, não seria capaz de ver se ele trouxesse, se ele não permitisse a quantidade de lances que ele tinha que analisar. Essa Bolívar falou uma outra coisa curiosa, eh, que talvez seja, não tem muito a ver com isso, mas só o comentário rápido. Ele achava que o jogadores de xadrez postal, não, não, como

é que era? A força é preserv..., a posição no ranking é mais ou menos preservada de acordo com a velocidade com que você joga as partidas. Então, os melhores jogadores do mundo em rápidas estão próximo entre os melhores jogadores do mundo de partidas não tão rápidas e partidas clássicas. E ele colocava e ele tinha essa teoria, mas ele colocou de outra forma. ele falou que os jogadores de xadrez postal, eles mantinham a posição no ranking aproximada com a posição dos jogadores eh clássicos, de clássicos. Eu fiz análise disso e tava tudo furado, mas não, eu fiz análise estatística, eu fiz eh aqueles padrões de pontos para saber, coloquei em eixo cartesiano, verifiquei tudo direitinho com ajuda de computadores, escrevi um programa para ver isso e tava furado. Mas era uma ideia muito interessante que a força se preserva independente da quantidade de tempo. Agora eu posso fazer o seguinte, a força se preserva se você jogar um computador de xadrez jogando contra outro computador de xadrez, a força se preserva com 3 minutos, depois com 5 minutos, depois com 10 minutos para cada um? Pergunta curiosa, né? –

(LPA) Eh, a princípio, a sim.

(CC02) Eu também acho que sim, mas eu não vi ninguém fazer esse teste.

(LPA) Aham. Não, curioso. É realmente curioso. E aí, eh, me vem uma ideia, por exemplo, eh a gente pode, a gente vai chegar numa inteligência geral, uma inteligência artificial geral.

(CC02) Eu não acredito porque, como eu falei, na minha concepção, inteligência artificial é um sistema especialista, ele resolve determinado problema. Por que de resolver esse problema e quando resolver esse problema é uma coisa mais humana?

(LPA) O Demis Hassabis, né, ele fala que muitas experimentações das neurociências são aproveitadas para ciência da computação. Você acha que a desenvolver esses algoritmos eh depende dessas descobertas da neurociência? Como que o cérebro trabalha, como que ele faz para se desenvolver ou não? Eh, seria realmente resolver problemas?

(CC02) deixa eu pensar um pouco. É, se eu entendo direito, você tá dizendo que à medida que você descobre como o cérebro funciona, você aplica isso em programas de computador?

(LPA) Eh a ideia das redes neurais acho que partiu um pouco disso, ou não?

(CC02) Não, quer dizer, partiu do conceito, você basicamente você pontua eh você cria vários vértices e em cada vértice você computa uma certa característica. E essa característica você exporta pros vértices que estão ao seu lado. Então elas elas propagam. O conceito é semelhante ao conceito neuronal, que o conhecimento é propagado através da da das células cerebrais. Então, nesse sentido, são iguais, só que um é matemática, eh, e o outro é um é digital e outro analógico. Não sei nem se é analógico ou não, mas me parece que é analógico.

(LPA) É, são bem as ideias do Nicolelis também, não é?

(CC02) Isso não, não sei se são bem as ideias dele, mas eu vi várias coisas dele, mas essa ideia eu já tinha antes também. Quando eu aprendi redes neurais, eu entendi dessa forma.

(LPA) Uhum. Tá. Então, então, deixa eu fazer uma pergunta. eh que eh segundo palavras do Stephen Hawking em 2014, que o desenvolvimento da inteligência artificial completa pode significar o fim da raça humana. O que que eh você acha desse tipo de declaração sobre os perigos da inteligência artificial?

(CC02) Acomodar. As pessoas se acomodarem e deixar as tarefas nobres para aquelas máquinas. Eh, a máquina do tempo de de HG Wells, para mim é aquele conceito, as se você deixar as pessoas eh eu acho que as pessoas são atraídas pelo por aquelas tarefas que exigem um menor quantidade de esforço possível. É uma tendência, não para todas as pessoas, claro, mas grande parte das pessoas tem uma tendência a escolher tarefas que você não precisa se esforçar muito. Se você der, eu acho que o sentido que ele quis dizer é o seguinte, se você permitir isso pro ser humano, ele vai adotar isso e aí sim, conceitualmente ele pode causar uma derrocada da raça humana.

(LPA) Uhum. Você assistiu aquele filme 2001 Uma Odisseia no espaço?

(CC02) Sim, li o livro, livro umas três vezes e o filme umas duas vezes só.

(LPA) O, e tem um tem um momento do filme como o algoritmo Hal 9000, ele tá programado para eh realizar a o objetivo, né, da daquela missão espacial. E ele percebe que os astronautas acham que ele tem um defeito e que querem desligá-lo. Então, para ele cumprir a missão da qual ele foi programado, ele resolve eliminar os astronautas.

(CC02) Isso. Eh, essa é a questão de prioridade.

(LPA) Ou seja, é o ser humano perdendo o controle da sua ferramenta.

(CC02) Sim. É, nesse ponto o Arthur Clark quando escreveu o livro, ele não acompanhou por algum motivo, acho que é o motivo é literário mesmo, ele não quis entrar nas três regras dos robôs de Asimov, que a primeira regra é jamais fazer mal ao ser humano. Eh, aliás, tem um dos livros dele tem uma história do robzinho que vai, só floreando um pouquinho, robzinho que está em Mercúrio, não sei se você leu esse livro?

(LPA) Não, não.

(CC02) ele tinha a tarefa de cavar de mineração, só que quando ele ia para o local de mineração, de mineração, que tinha uma pessoa presa lá embaixo, se eu não me engano, ele ia até lá Só que pegava muito sol, ele tinha que voltar. Então ele ia e voltava e voltava e voltava. Ele ia quando ele percebia que ele ia morrer se continuasse e com isso matar o ser humano, porque o ser humano que estava preso lá não conseguia sobreviver, ele voltava. Então ele ficava entrando, indo e voltando. Então é a questão de prioridades. A prioridade dele era salvar um ser humano. Só que ele não podia fazer isso sem se sacrificar. Só que se ele se sacrificasse, matava o ser humano do mesmo jeito. E no caso do Hal, eh, convenientemente, eu acho que foi pro conceito do livro, o conceito do livro diz que, eh, você não tem a regra dos três robôs e as três regras dos robôs. Então, o que o Clark quis fazer é um livro aonde houve um erro de programação sobre quais são as prioridades que devem ter para a missão. Ou talvez, eu não me lembro se foi intencional ou não, não me lembro se em algum momento, se em 2010 ou 2061, os dois livros que seguiram, se eles comentam. E eu tenho uma ideia vaga de que o erro foi o a prioridade foi alterada, acho que 2010 que a prioridade foi alterada para a missão e não para as pessoas. Acho que eu já viajei na maionese também.

(LPA) E o Hal 9000 não contava com a inteligência humana em desligá-lo, né, em conseguir ... Eu achei muito e assim ...

(CC02) Exatamente. Exatamente. Ele não previa, não previa todas as possibilidades.

(LPA) E o que que deve ser relevante pro desenvolvimento das inteligências artificiais daqui para frente? Eh, que tipo de pesquisas devem acontecer e se o xadrez pode ser útil para esse tipo de pesquisa?

(CC02) E eu acho que xadrez é uma, são blocos diferentes. Vamos pegar a inteligência artificial, o termo tão utilizado, vamos pegar o chatgpt ou outra ferramenta qualquer. Ela é uma ferramenta de auxílio e de estruturação de raciocínio em alguns casos. Eu fiz várias

buscas na internet, várias buscas no Google e a busca da a resposta da inteligência artificial do Google normalmente já é o suficiente para mim. Ou seja, eh, na busca de uma coisa especialista, uma coisa bastante precisa, bastante bem definida, ela responde muito bem. Eh, já eu acho que o ser humano ele é quem vai fazer, ele vai é quem vai comandar as pesquisas, quais são as pesquisas que devem ser feitas eh, como implementar ou como reagir a essas pesquisas. Então essa é uma coisa. A outra coisa é o xadrez. O xadrez é um, na minha maneira de ver, ah, o estudo de inteligência artificial do xadrez. É um sistema especialista também, mas é um sistema especialista que procura entender como que você pensa problemas complexos, mas com regras bem definidas. Já o mundo não tem regras tão bem definidas. É isso, isso que eu queria chegar desde o começo. Isso que eu tava procurando, tentando chegar. Como a vida humana não, a vida não tem regras bem definidas, você não tem como, eu não vejo como se chegar a um sistema especialista que resolva, senão você vai chegar e vai dizer que a solução da vida é número 42, né? Você sabe o que que é isso, né? –

(LPA) Uhum. Não.

(CC02) Eh, isso aí é o mochileiro, guia do mochileiro das galáxias, que uma máquina foi construída para responder a grande questão do universo. Qual é a grande qual é a razão da vida? A máquina pensou por 100 milhões de anos e respondeu: 42. Aí eles montaram outra máquina para saber qual é a pergunta cuja resposta era 42 que ele respondeu. É um besterol completo aquele livro lá, mas é muito divertido.

(LPA) Ah, já ouvi, já li em algum lugar sobre isso também. E você queria, gostaria de acrescentar alguma coisa? (CC02), sobre a questão xadrez, inteligência artificial ?

(CC02) Eu acho que se não fosse divertido trabalhar com xadrez, se não fosse um jogo tão bem definido e com tantas possibilidades, você não tem veja, a quantidade de opções que você tem na vida são infinitas. Xadrez são finitas, porém eh, eu diria contáveis, mas quase infinitas. Existe uma diferença grande. Você pode contar a quantidade, você pode enumerar quantas possibilidades de jogo, de de tabuleiro de posições diferentes você vai ter. Eh, e pode traçar um vértice entre cada posição, de uma posição para outra, para verificar qual que é o melhor desempenho, qual o melhor jogo. Então, o que me parece é que o xadrez sim pode ser uma, uma maneira de estudar, como é que eu vou falar? Um ambiente controlado. Isso que eu queria. Um ambiente controlado para você fazer experimento sobre o infinito incontável. Cara, até parece que eu sabia o que eu tava

falando. Mas é mais ou menos isso mesmo, porque você é, é um local que você consegue analisar o infinito contável para tentar chegar ao infinito incontável.

(LPA) Belíssima frase essa, hehehehe.

(CC02) Cara, até parece que fui eu que falei. Falei, soou bem, né?

(LPA) (CC02) Tô satisfeito com o que a gente conversou aqui. Vou parar a gravação.

(CC02) Joia.

APÊNDICE E – ENTREVISTA NC01

Íntegra da Entrevista com a neurocientista NC01 (27/08/2025)

(LPA) Pronto, agora sim. Então, obrigado pelo pela aceitação do convite, professora. Eh, fico muito contente e orgulhoso.

(NC01) eu é que fico. Obrigada a você.

(LPA) Muito bem, professora, são perguntas muito gerais e fique à vontade para se aprofundar, no que for conveniente e de repente e que possa contribuir para a nossa pesquisa de acordo com o que eu relatei; mas eu queria começar com uma pergunta que que a gente comentou ah recentemente, que é sobre uma frase do professor Nicolelis, né, que ele fala se IA aqui a IA não é nem inteligente nem artificial. Gostaria de um comentário teu, do seu entendimento, do que que ele quis dizer.

(NC01) Eh, então acho que um tanto do que que eu interpretei, né, baseado também na minha experiência. Primeiro acho que a gente tem que eh entender eh o conceito de inteligência, né, que é algo que é sempre debatido, ou seja, uma habilidade numa tarefa específica, né? A pessoa pode ter uma superhabilidade numa tarefa, mas ela não tem habilidades, outras habilidades, inclusive compatíveis com a vida em sociedade, em grupo, né? comunidade, etc. Eh, então, o fato de haver, e aí eu já antecipo um pouco, né, do do que eu entendo por IA e eu tô ainda aprendendo, né me instruindo sobre o tema, eu vejo como um meio de reunião de organização de informações, né, de muitas eh dados que foram gerados, né? Então eu entendo a IA como uma ferramenta de organização daquilo que já foi produzido, né? Então, eh, e eu não vejo como artificial, porque de fato um componente da criação, né, que aí deriva de de muita coisa, de experiência, de contexto cultural histórico. Eh, eh, eu vejo com muita restrição a possibilidade de ser inteligente de fato no sentido da inovação, criatividade, né, e da, então, na verdade, é uma ferramenta que reúne coisas que já foram produzidas, né? A minha facilitará o entendimento em, especialmente em algumas áreas, sem dúvida, né, inclusive antecipando coisas. Eu vivi uma situação, né, pessoal, desse processo que eu contei um pouco mais cedo, que eu tive uma complicação de uma cirurgia que eu fiz e o chat GPT diagnosticou 15 dias antes dos médicos com uma iniciativa feita pela minha filha que jogou no chat GPT os sintomas.

(LPA) Olha. Uhum. –

(NC01) E aí, obviamente, então acho que assim, eh, pensando na medicina, né, a partir desse exemplo, eu acho que os profissionais, eu não vejo problema nenhum e os profissionais admitirem que não vão conseguir lembrar ou ter conhecimentos todos do funcionamento do corpo, que uma ferramenta como essa pode antecipar diagnósticos, pode. Então acho que a gente vai até no processo de formação, né, de que nós somos professores, a gente tem que aprender a como lidar com isso. A gente já viu uma realidade que eh alguns estudantes, a gente tá falando na sala de aula, eles estão checando se o que a gente tá falando tá certo, né? A gente já vive essa realidade. Eh, não é de agora, né? A própria ferramenta de busca da Google já proporcionava isso, né? Então, eh, a gente vai ter que aprender a lidar, mas eu concordo com Nicolelis, né, de que não é inteligência, uma vez que inteligência ela é multifatorial, ela é deriva de um sistema complexo, né, enfim. E e não é artificial, uma vez que ela simplesmente reúne aquilo que foi humanamente produzido, né? Ou seja, a partir da visão dos humanos mesmo, né? Então eu não vejo uma autonomia, né, dessas ferramentas sem o fator humano por trás, né? Uma eu não vejo de fato. E o termo de fato ele não é legal, né? Eu vi que no título do seu trabalho você usa inteligência humana, né? Eh, e eu até para contrapor o artificial tenho feito o debate inteligência natural, né? Porque se for eh, é claro, a gente ir pensando, né, na no atual no atual debate absolutamente necessário, existem outras inteligências, né? Se a gente for pensar inclusive plantas que têm ciclos, né, circadianos, biológicos, né, ou seja, tem aí. E a gente não pode ignorar isso,

(LPA) Tem um livro que eu ali que é da Lúcia Santaella, não sei se você conhece, e ela fala que a gente tem uma tendência

(NC01) né? Hum. Eu já ouvi falar dela assim.

(LPA) muito a antropomorfizar toda essa –

(NC01) Isso,

(LPA) questão. Então, ela nas nos livros dela, ela comenta que até uma máquina de escrever pode ser inteligente. Uhum

(NC01) isso. Exatamente. É, é muito louco você passar, né, essa coisa para objetos, né? Mas inclusive em algumas culturas, né, culturas indígenas, eles eh a forma como eles

lidam, né, com o ambiente, com o entorno, a pedra tem representa inclusive ancestrais deles, os rios, né, ou seja, tem um conhecimento, uma inteligência, né, por trás dessa forma como lidar, né, que eu tava até eu tava debatendo esses dias, né, a gente tá aí fazendo, assinando acordos, protocolos e termos. Vai ali na comunidade indígena e pergunta como é que eles preservam, né? O porquê 5% da população mundial preserva 90% da biodiversidade, né? Isso é inteligência, isso, né? Ou seja, e os caras não tão com recursos nenhum tecnológicos na forma como a gente entende tecnológico, né? Porque tecnologia não necessariamente tem que resultar num equipamento, né numa ferramenta em si, né? Tem uma eh, então tem várias formas. Eu costumo dizer, Léo, se você já me escutou dizendo isso, né? me perdoe, mas assim, uma das pessoas mais inteligentes que eu conheci foi o meu avô paterno, que era trabalhador rural, aposentou com salário mínimo, morava numa casa que metade era de alvenaria, metade de pau a pique. e analfabeto, né? Ele não sabia sequer desenhar o nome, né? E ele era de uma inteligência assim na observação, não só de observar. Então, às vezes ele tava lá fumando o cigarrinho de palha dele, ele falava, olhava assim pro horizonte e falava: "Lá paraas 3 da tarde vai chover". Sem ferramenta nenhuma. Dava 3 da tarde, chovia. Ou seja, tem uma percepção ali, um conhecimento, uma inteligência que muito possivelmente tem relação com o vento, o canto dos pássaros, né, a incidência do sol. Ele não precisava de ferramenta nenhuma para fazer essa previsão.

(LPA) Provavelmente a sequência de perguntas que eu elaborei vai chegar aí, eu acho que sim. Vamos ver.

(NC01) né? Vai chegar aí, tá? Uhum.

(LPA) Eu começo assim. Eh, quais seriam os fatores que diferenciam máquinas e humanos no processo de aprendizagem?

(NC01) Nossa, muitas assim, eu sou supercrítica, a visão mecanicista que eh faz a associação do corpo com máquina, né? é uma visão eh não só reducionista, mas que dialoga com o desenho inteligente, ou seja, é algo que foi projetado para funcionar daquela forma. Como a gente tem uma perspectiva evolucionista, né, a gente sabe que inclusive somos frutos da seleção natural, do acaso, né? Então, eh, eu absolutamente discordo dessa analogia de qualquer parte do corpo humano como máquina, porque a gente o próprio processo de aprendizado são tantos os processos, os mecanismos que estão acontecendo, que tem um nível de aleatoriedade imenso, né? Os caminhos, os

resultados eles são absolutamente podem ser distintos, recebendo o mesmo estímulo. Você pensar que se você dar um estímulo, vai dar sempre o mesmo resultado é um equívoco, porque vai depender de uma série de fatores que da idade, se a pessoa tá bem, emocionalmente bem, do contexto cultural, se ela dormiu bem, se ela comeu bem, enfim, depende de uma série de fatores, né? Então, a gente chegou, não sei se a gente chegou a discutir, né? Ah, tem um livro chamado Dialética da Biologia, que é do Levontin, né, e do Levins, que fala um pouco desse diálogo, né, entre sociedade, né, que tá aí colocada no programa e os fenômenos, os processos biológicos, né? A gente não pode separar essas duas dimensões e não dá para hierarquizá-las. também, né? Ou seja, não tem uma que é mais as duas tão, né? Eu gosto sempre de usar um jargão que eu aprendi com o meu companheiro. Somos 100% cultura, 100% biologia, né? Então, pode falar.

(LPA) Eh, não, eu acho que encaixa com uma pergunta que eu tinha posto mais pro fim. Eu vou antecipar que eh se a inteligência artificial para aprender com o ambiente, ela necessita então de um corpo físico, ou seja, ter sensações para que se forme essa experiência, esse conhecimento?

(NC01) A inteligência artificial tendo sensações.

(LPA) Isso, porque o Antônio Damásio, ele fala justamente isso, que a inteligência artificial não tem um corpo que recebe a

(NC01) É isso.

(LPA) Mas que os cientistas da computação tão acoplando, por exemplo, uma câmera para ter a visão ou um, né, um fone, um microfone.

(NC01) Olha, eu não entendo nada de algoritmo e de mas hoje, por exemplo, eu dei uma aula sobre a audição, gustação e olfação. Eu juntei as três modalidades, né? E eu tava discutindo com os estudantes, né, a questão da individualidade da percepção daquilo. Então, eu posso receber o mesmo estímulo. E as percepções e as interpretações, elas são absolutamente distintas e baseadas na nossa história, na nossa memória, né? Então, eu usei o exemplo do coentro. o coentro que as muita gente sente gosto de sabão. Aí um aluno virou e falou assim: "Ah, mas já tem um gene que as pessoas têm, pessoas que têm esse gene eh eh tem mais chance de sentir gosto de sabão." E aí eu provoquei, eu falei assim: "Mas e se uma criança cresce escutando o pai e a mãe que coentro tem gosto de sabão?" Ela pode inclusive não experimentar porque ela acha que vai ter gosto de sabão.

E um belo dia ela pode comer e não vai sentir gosto de sabão. ou ela sente, porque a construção mnemônica que é depende um tanto da subjetividade, né, individual e não à toa, a gente tem aí discussões seculares sobre corpo e mente, né, essa separação, né, que não existe de fato na minha na minha visão, né? Então, assim, além das individualidades todas, os processos eles são muitíssimo distintos. Como é que você vai programar um algoritmo que vai simular? Vamos pensar quantos quantos bilhões de habitantes a gente tem no mundo

(LPA) uns 8 bilhões.

(NC01) atualmente? 8 Bi 8 bilhões de pessoas. E cada uma dessas pessoas podem ter respostas diferentes para para estímulos que sejam muito parecidos, né? Então, e aí tem a questão da volição, que é uma grande vontade de eu ter alguma coisa que depende de vários processos, não só biológicos, né, mas enfim. Então, eu de fato não acredito que eh algoritmos e que máquinas simularão tudo aquilo que biologicamente, culturalmente é construído, até porque a gente tá o tempo todo mudando, né? Eh, eles estão simulando algo que já foi. Inclusive, se a gente de fato acredita na criatividade, na imaginação, na inventividade, que é algo que nos diferencia inclusive de outros primatas, talvez nem tanto, né? Porque macacos também constroem ferramentas, eh, e passarinhos, pássaros também, né? Tem os pássaros que usam ferramentas, eh, e é impossível porque eles vão estar sempre olhando pro passado, né, e não para esse futuro possível, que é algo sensacional assim, né, da mente humana.

(LPA) interessante é que eh tá concatenando as ideias, porque a minha próxima pergunta e a eh –

(NC01) Eu tô eu tô antecipando respostas. É isso?

(LPA) não tem essa coerência lógica, né, desse raciocínio todo que a gente tá fazendo, que é o objetivo, afinal de contas

(NC01) Sim,

(LPA) né? E eu ia, a pergunta justamente é essa, né? Simular o pensamento humano é possível por uma inteligência artificial. As máquinas poderiam simular a consciência.

(NC01) Eh, então qualquer simulação vai ser a partir daquilo que é fornecido a elas. E o fornecido a elas é a partir de conhecimento que já foi. Então, se eu pensar, né, no próprio

processo evolutivo, né, e eu acredito que a gente continua, vai continuar esse processo, é que a gente não vai ter tempo de ver, até pensando no tempo histórico, né, do estabelecimento da ciência enquanto instituição, digamos, né, ou religião, seja lá o que for, a gente não vai ter tempo de ver, mas certamente isso tá continua acontecendo. Eh, então esse processo de mudança que vai desde a intimidade do gene, né, e até aspectos mais gerais da cultura, da história, da organização social que a gente escolhe, né, e tudo muda, né? Eh, a gente vem de uma cultura, eh, por exemplo, né, da coisa da imaginação, eh, majoritariamente politeísta para monoteístas, né? Ou seja, isso é uma mudança radical e a gente tem eh consequências até hoje, né? E aí a interpretação imaginativa, né? Porque alguém já viu Deus? Eu costumo dizer, né, do alto do meu ateísmo que Deus e a outras coisas mais eh na visão materialista, digamos, é a maior, é a fake news que mais fez sucesso, hein? Porque é uma história inventada na minha concepção, né? E que as pessoas de fato acreditam e seguem e defendem e se organizam em torno disso. Então, a gente fica fazendo toda essa discussão sobre fake news, né? Mas ela existe há muito tempo, né? Há muitos séculos. Então, eu não acho que eu discuto isso, que vai conseguir simular. Eu, as pessoas têm falado que estão assustadas com a IA. né? Então eu não tenho experiência, para ser muito honesta, assim para dizer desse susto de colocar o componente humano, tal. E mas assim, eu reconheço os textos dos estudantes quando é escrito por IA, porque eu sei como é que eles escrevem, né, pela minha experiência profissional. Eh, tem alguns jargões, tem alguns vícios que são identificáveis para quem é leitor de literatura. e ler outras coisas, a gente reconhece, né? Eu sei que as versões pagas estão e é disso que se trata, né? É vender produto. No final as versões pagas estão melhores, né? Mas eu acho que aí tem toda uma discussão moral e ética também por trás, né, do uso, porque é quase como como eu encaro a IA como uma ferramenta, né? Eu não acho que é um é algo que é o fim em si mesmo. Então, eh, eu tenho discutido com os alunos, vocês estão enganando a si próprios, é como se tivessem colando. Se eu peço para vocês produzirem o texto, vocês vão lá e pedem para esse organizador de dados gigantesco, porque é gigantesco, né? Produzir um texto para você. Vocês estão enganando o professor ou se tá enganando a si próprio,

(LPA) Olha, olha como é interessante, né? Eh, também concatenando nessa pergunta, né? Eu não sei se esse discurso eu consegui pensar em tudo isso, mas me parece que tá sendo

(NC01) Sim,

(LPA) concatenado. A pergunta seria: pode a inteligência artificial se tornar social, interagir umas com as outras? Como as neurociências interpretam esse tipo de questão?

(NC01) Então eu imagine, né, Léo, eu sou da das neurociências, mas eu sou professora de anatomia, eu sou fisioterapeuta e eu sou corpo. Eu sou corpo e para mim não existe. E a gente viveu um período agora que eu acho que responde muita dessa questão que foi a pandemia. E o fato da gente ter que fazer tudo online. Não à toa, as pessoas, a gente tava conversando mais cedo, começaram a adoecer. Então assim, tem componentes da do próprio sistema sensorial, o cheiro, o toque, a textura, o ouvir, que como é que você vai simular um beijo de língua com o IA? né? Ou seja, você teria, eu sempre me lembro, né, de quando eu fui, eu fui fazer uma disciplina na Poli, porque era uma disciplina que falava da do do componente elétrico, elétrico de moto neurônio. Era uma coisa bem da elétrica, era na engenharia elétrica. E aí eu fui fazer essa disciplina totalmente outsider, assim, né? Não entendia nada de só lá no meio dos engenheiros. O professor, que era o André Fábio Con, eh, me sugeriu trancar a disciplina e eu teima, falei: "Não, eu vou até o fim". E aí eu fiz amizade com com um físico que tava fazendo hã uma modelagem de um motoneurônio. Na época, isso lá no início dos anos 2000, não existia computador com memória para ele simular todos os processos que aconteciam em um motoneurônio, porque ele tinha que simular tudo, todos os canais iônicos, todos os receptores, todas as organelas, a condução do estímulo elétrico. E isso sem pensar na sinapses. o motoneurônio, né, ele pode ter no mínimo 1.500 contatos sinápticos diretos, sem contar rede eh neural. E isso a gente tá pensando na lógica só do neurônio. Quando a gente bota células gliais, que são outros bilhões de células, também se comunicando de uma forma que não é tão facilmente registrável quanto o neurônio é, que tem o potencial elétrico, né? Ou seja, por isso que eu falei, eu gosto, o que melhor simula para mim o funcionamento do sistema nervoso e o resultado da interpretação de um estímulo que a gente recebe, eu uso sempre a imagem do efeito borboleta, a teoria do caos. teoria do caos, mesmo assim, a gente não você achar que o mesmo estímulo vai ter sempre o mesmo efeito é um equívoco. Então eu não acho que vai se relacionar porque os próprios relacionamentos são dinâmicos, né? Ou seja, tem saturação dos sistemas, né? A gente se dá bem com uma pessoa, daí a pouco já não tá mais dando certo por várias razões, porque tem um componente da emoção que como é que você vai simular a emoção?

(LPA) Era justamente minha próxima pergunta. Eu ia falar assim, -

(NC01) Não é só rir. Ah, é. Olha só. Tá vendo?

(LPA) um dos próximos passos a ser dado pela inteligência artificial é de colocar emoções no processo de tomar decisões.

(NC01) Então, eu acho que assim, pensando no objetivo da inteligência artificial, isso eu tô pensando nas ferramentas que eu conheço, né? O objetivo dos caras que estão por trás dessa tecnologia é vender coisas. Então eles podem até conseguir deflagrar emoções nas pessoas, assim como a propaganda da Dorian na TV deflagrava, só que não era numa não era tão exponencial quanto é agora, né? Então eles vão lidar com emoções como uma ferramenta de marketing, inclusive para poder fazer com que ela aquela pessoa se atraia por aquilo. E tem um componente que é do ser humano também, que a gente se vicia nas coisas, né? A gente se vicia naquelas coisas que exatamente são as que nos geram emoções, sejam emoções pro bem ou pro mal. Eu gosto sempre de falar pros alunos por será que o corpo ferido atrai a atenção das pessoas, né, de muita gente. As pessoas querem ir, tem uma pessoa passando mal, tem uma pessoa que se acidentou, as pessoas querem ver, né? Porque isso lida com as emoções, né? Isso mexe com as emoções. E os caras sabem disso, né? Eu acho que naquele documentário lá que retrata a história do Zuckerberg, né, do Facebook, né, eu acho que eh é um é um bom retrato de como ã como essas estão pensando o uso dessa tecnologia. Eu me lembro, Léo, não sei se eu comentei, eu comentei a história do big data, né, do neuroscience. Você me escutou comentando isso?

(LPA) Sim, acho que sim. no finalzinho.

(NC01) Porque assim, uma coisa é usar uma ferramenta que vai conseguir juntar, por exemplo, eh, os resultados de dezenas de milhares, centenas de milhares de estudos e conseguir entregar algo que seja uma síntese, né, uma categorização, que eu acho que isso é uma coisa boa. Outra coisa é pensar no uso dessa tecnologia para vender, para viciar, para eh alienar as pessoas, que é isso que tá acontecendo. As pessoas estão se alienando. E eu trabalho numa perspectiva de desalienação do próprio corpo. Eu faço experimentações com os estudantes para eles sentirem no corpo as sensações para depois pensar nos mecanismos e na células e no que tá acontecendo, né? Para que se a gente tenha consciência, então a gente tá com problema, as pessoas estão se tão sem consciência

de si próprios, né? Então eu acho a gente, eu não sei o que vai acontecer, né, com crianças, idosos. Eu sempre acredito que pegando de novo a experiência da pandemia, chega um momento que satura e as pessoas querem buscar o contato. Hoje um aluno me falou que desinstalou do celular todas as redes sociais e que ele sofreu muito nas primeiras semanas, primeiros 15 dias, mas que agora ele não sente mais falta.

(LPA) Ótimo. –

(NC01) Então tem um retorno aí também, um racionalizar sobre também o que significa isso, né, que é importante, a gente vai aprendendo. Ah

(LPA) Uhum. Então, e quanto a metáforas e analogias, elas são processos de simulação mental exclusivos dos seres humanos?

(NC01)) ah, então, nossa, essa pergunta a gente vai pra área da filosofia, né?

(LPA) É.

(NC01) Vamos filosofar. É. Eh, então a inventividade, a capacidade imaginativa, se eu for para olhar para as estruturas, né, e comparar outras espécies animais, né, se for pegar os primatas, a gente de fato tem uma estrutura que é eh filogeneticamente mais recente, que é o neocórtex. E é exatamente no neocórtex que esses processos cognitivos, intelectuais eles ocorrem. Então essa capacidade antecipatória, imaginativa, criativa, de fato, a gente até cola... , a gente identifica alguns sinais em outras espécies, né? Mas muitas das vezes tem relação com sobrevivência. Ponto, né? a gente. Então assim, se a gente pensar nas artes e na capacidade infinita de criação, porque veja, as artes traz uma e as artes são essenciais, né, na formação do indivíduo e e assim e possibilidades inúmeras, né, de criação. Eu sei que a IA tá criando música, inclusive, né? escutei na rádio ontem que eh estão atribuindo a artistas que não gravaram o disco, a obra, né? Aí cai de novo na discussão ética e na questão financeira de venda, porque as pessoas estão fazendo isso para vender. Eh, mas eu não sei o quanto é criativo, inovador. Eu não, eu não vejo esse potencial de inovação. E eu assim escutando, né, lendo superficialmente, os próprios criadores de algumas IAS falam disso, né, que esse esse eh capacidade de criação, ela é muito do Homo sapiens, né, então de inventividade. Então, pode até ser, mas eu, como eu disse antes, eu não vejo um uma autonomia total de máquinas e até porque depende de quais são os recursos, né? Os recursos que eles precisam são finitos, os tais das terras raras com seus metais são finitos né?

(LPA) E as fontes de energia, né, para alimentar ...

(NC01) As fontes de energia, né? Enfim. Eu eu não sou uma grande eu já fui acusada tecnofóbica. Não sei se eu já te contei isso.

(LPA) Não ... ludista?

(NC01) não, porque eu tava discutindo a história da plataformização e eu tava discutindo numa perspectiva meio de negar, sabe? Eu tava muito revoltada com essa coisa dessa transferência da vida para pro virtual, né? E aí a pessoa falou: "Você tá me parecendo tecnofóbica?" Eu falei: "Não". Não, assim, eu tô também aprendendo, né? Mas eu não vejo aqui do meu olhar de das neurociências substituição de todos os processos. Não vejo.

(LPA) Uhum. Olha, eh, eu acho que foi bem interessante, (NC01), e até eu pensei numa pergunta, até tinha anotado antes, que não tava aqui no meu roteiro, mas é que eu acho que faltou mesmo, porque eh até pela pelas suas palavras, né, a gente percebe que muita coisa das neurociências foram aproveitadas pelo pessoal da tecnologia, né? Mas teve alguma coisa da tecnologia que ajudou as neurociências?

(NC01) Ah, sem dúvida, né? Assim, primeiro eu acho que a própria tecnologia, o conceito de tecnologia acabe também em discussão, né? Porque não necessariamente eh a o estudo da técnica, a tecnologia, ela resulta num produto, num artifício, etc. Mas certamente todas as ferramentas de as análises, né, os métodos, muitos deles inclusive nos levaram à possibilidade de reafirmar aquilo que Santiago Ramon y Cajal, por exemplo, colocou na teoria, né? Então ele tinha ali o seu microscópio monocular usando o embrião de galinha e ele dizia coisas que a gente só tá conseguindo visualizar, né, e comprovar hoje. Não à toa volta e meia tem lá menção Santiago Ramon y Cajal , né? Então, eh, certamente a gente não pode negar todos esses recursos, mas eu tenho consciência que muitas dessas medições desses recursos, eles não estão fechados em si. Ou seja, a gente tem um desafio muito grande para quem é da pesquisa básica, né, da pesquisa dura, de que se eu trabalho com uma molécula, com um mecanismo biológico, eh, eu, obviamente, a descrição desse mecanismo é importante, mas eu não posso negar todo o entorno, né, inclusive o próprio entorno celular. Você tem uma ideia, eu tenho trabalhado com um conceito, tem aquela coisa, né? A unidade funcional do sistema nervoso é o neurônio, tá? Eu não gosto desse conceito, porque assim, o neurônio sozinho, o carinha não vai fazer nada, ele tem que ter as células gliais, ele tem que estar em contato com outro neurônio. Então eu gosto de

trabalhar com conceito de unidade funcional, ou seja, é um trabalho em grupo, né? Se esse grupo não tiver funcionando, nada funciona. E quando a gente vai pro corpo humano é um tanto disso, né? Então, do que adianta eu simular, eu fazer a modelagem de um moto neurônio, se eu não tô pensando em toda a integração que ele eh vai fazer, né? Então assim, certamente a tecnologia ajuda a responder perguntas, mas e eu acho que a IA poderia nos ajudar a integrar essas coisas, sabe? Então, alguém analisou uma proteína aqui e tal, é é isso que eu eh como eu vejo, né, de eh trazer os resultados, uma síntese daquilo que já existe, uma organização dos dados, etc. Eu vejo como um grande organizador potencial de dados, mas geração de coisas novas, tal, é só a partir do que já existe na minha cabeça, sabe? Mas enfim,

(LPA) Não ... Para mim tá ótimo, (NC01). Foi, para mim foi riquíssima essa. Eh, você gostaria de colocar, enfatizar alguma coisa, acrescentar –

(NC01) sim, tá? Acho que acho que não, acho (que talvez essa coisa do que eu falei de 100% biologia, 100% cultura e acreditar na dinâmica dos sistemas e dos processos, né? Porque tudo tá ainda acontecendo, não à toa alguns biólogos estão trazendo o lamarquismo pro debate, né? Porque eh obviamente não é da forma como o Lamarck descreveu, mas toda a área de epigenética tá um pouco resgatando aquilo que ele falava de adaptação, né, e de mudança estrutural a partir do contato com o ambiente. Ou seja, é tudo muito extremamente dinâmico assim, né? E ao mesmo tempo lento quando a gente pensa no tempo evolutivo, né, que a gente Então acho que isso a gente não pode reduzir, por isso não reduz a máquina. Sinto muito, mas não somos uma máquina, né?

(LPA) Ótimo. Tá ótimo. Vou parar a gravação.

APÊNDICE F – ENTREVISTA NC02

Íntegra da Entrevista com a neurocientista NC02 (17/09/2025)

LPA: Então é isso, dr., muito obrigado. Eu gostaria de começar com uma pergunta que muita gente tem feito e vem lá do de uma frase do professor Miguel Nicolelis que fala que a inteligência artificial não é nem inteligente nem artificial. Eu só gostaria que o senhor comentasse essa frase do Nicolelis.

IDV: Uhum. Eu gosto dele também. Então é opinião, né? em medicina ajuda muito, né, que nem eu mesmo, fiz cirurgia robótica. Ao invés de abrir, foram feitos seis furos. Depois eu mando para você. Eh, e ajuda muito porque você faz uma tomografia, aí você marca pontos, por exemplo, no crânio, né, crâniométricos, porque uma lesão grande, você abrindo o cérebro, você tá acostumado, você acha, mas lesão pequena tem como se fosse um GPS. que você faz uma tomografia, marca o crânio, marca onde está, e aí você vai com a agulha e vai mais para a direita, mais para a esquerda. É como se fosse quando aquelas aquela voz do GPS você está a 200 m, vire à esquerda 300, entendeu? Quer dizer, dentro da neurocirurgia ajuda demais a achar pequenas lesões. Uma outra coisa que ajuda, hoje em dia tem muito médico que atende com o chat GPT aberto, é por quê? Ele está dentro ali e já está sendo gravado. Ele se defende. Se uma pessoa disser que ele ofendeu, que ele não deu atenção, ele tem gravado. A pessoa no fim da consulta fala assim: "Eu queria um relatório". Ele já tem um relatório, né? Agora o que acontece é o seguinte, né? ainda é todos esses computadores eles estão sendo constantemente alimentados, né, pelos algoritmos, né? Eu ainda tava numa reunião sobre o chat GPT, um monte de médico jovem, eu tava mais ouvindo, né, vendo assim dose de remédios, né, conforme a idade e tal. Então o pessoal comentando tal e numa mesa de jantar, né, uma mesa grande, uma confraternização e todo mundo conversando e a moça da minha frente falando, tal, né? Aí uma bela hora ela, uma pessoa falou assim: "É, mas de vez em quando ele erra". Aí a moça apontou, falou assim: "Por isso que o senhor precisar eu vou nele.". Apontou para mim. Eu falei: "Tá certo, não que eu não erre, né? Não que não erre. Claro que imagina a gente é humano, né? Mas tô dizendo, a probabilidade é menor de fazer um erro grande.

LPA: Sim. Então, o senhor discorda da frase do Nicolelis?

IDV: Não, eu acho que eh principalmente em medicina é muito válido, né? É muito válido. O cálculo pelo computador você vai com precisão. É um GPS, né? Mas orienta bastante.

LPA: É uma espécie de inteligência?

IDV: É uma espécie de inteligência que nos ajuda muito, dose de remédio. Eu estou conversando com você, você fala: "Pois é, e quantos anos você tem? Qual sua idade?" Eu já sei o cálculo da medicação, né? daria para fazer na calculadora, mas já tá tudo agendado ali. Eh, que nem o profissional, não é um neurologista, mas tá lá enxaqueca. Tô com dor de cabeça. Tipo assim, ele já viu que ele é uma enxaqueca. Que que você costuma usar? Ele pode não ser especialista, mas ele já se orienta, Não é um especialista tal, mas é um orientador, né? –

LPA: Sim, entendi. é que a frase dele ficou bastante polêmica, né? Então, eu gostei de tocar nesse assunto.

(NC02) É, não, eu gosto dele. Eu gosto dele muito dele.

(LPA) Uhum. E quais seriam os fatores que diferenciam as máquinas e os humanos nesse processo de aprendizagem?

(NC02) O, eu tenho um slide de quando eu dava aula que diz que sem, é um livro do Chalita, sem afeto não existe educação.

(LPA) Ó, legal. Muito bom.

(NC02) porque você pode ter informações, né? O computador que nem o chat GPT ali e tal, ele vai te dando uma série de informações. Acontece que principalmente a figura do professor, professor é que junta essas informações. Eu sempre quando eu dava a olha eu falava assim: "Até eu entender que um monte de bolinha, pode ser 10 bolinhas, esses esse 10 eu tive que aprender. Alguém me ensinou que aquilo juntando é 10, aquilo juntando é 100. Eu ia ficar vendo um monte de bolinha, eu não ia chegar nesse número e a a e, né, eh eh a informação, porque dentro do nosso cérebro tem uma área que ela é maravilhosa que chama área de associação, né? O conhecimento não é só isso, mais aquilo, mais aquilo e mais aquilo. O conhecimento é nessa área de associação acontece o milagre da evolução humana que chama o então. Então aquilo isso mais isso mais isso não é isso mais isso mais isso dá uma outra coisa. É o então é quando a criança fala "Nossa, nó

(LPA) É, eu digo que é quando acende uma luzinha na cabeça.

(NC02) então isso é a conclusão. É porque na área aquilo foi associado e se transformou numa outra coisa. Esse que é o milagre você, o mundo continuar evoluindo, porque as informações elas vão chegando, mas quem junta as informações é professor. Até eu entender que na neurocirurgia, para fazer uma cirurgia, eu tô vendo todas as ferramentas, mas a sequência, como é que eu vou chegar nisso aqui? Eu preciso ter a sequência para saber onde eu quero chegar. O Dr. Zman falava sempre para mim: "Quem não sabe o que procura não percebe quando encontra". São frases feitas, mas frases interessantes, né? Então não pode ...

(LPA) Uhum. Sim, inclusive. Ah, desculpa. Não, eu ia complementar com uma outra pergunta que tem referência do Antônio Damázio, que ele diz que a inteligência artificial para aprender com o ambiente, será que ela necessita de um corpo físico? Porque ele fala muito que o ser humano tem as sensações, eh, para que ele forme essa experiência e esse conhecimento?

(NC02) Eu acho que precisamos do corpo físico ainda, né? Eh, principalmente, por exemplo, numa classe tão heterogêneas são as classes, né, que você tentar dar uma aula sem saber que plateia que você vai pegar, vai ser uma aula inútil, né? Alguns vão pegar, mas a maioria não vai pegar, né? Eu tô falando de experiência de 14 anos, né? Você tem que chegar na classe e fazer algumas perguntas para se situar onde é que eles estão, né? Senão você fala muita coisa, mas você não tem utilidade, né?

(LPA) E acaba acontecendo essas associações também, não é?

(NC02) É, eu acho que fica muito frio, né? Porque eh as nossas emoções, eu lembro que quando eu ainda falava há muito tempo, eu tô usando umas frases da época, é como se fosse quando você tá trabalhando no antigo Word. Lembra naquele antigo Word? Você tem que ir fazendo as coisas e tem que ir salvando, né? Salvar. Salvar. Senão um pico de energia ou alguma coisa, quando você liga outra vez não tem mais nada lá, né? As nossas emoções são assim. As emoções é que fixam a nossa memória, tá?

(LPA) Não, muito interessante. Acho que nós vamos chegar lá na minha pergunta. Eu primeiro queria saber assim, a simular o pensamento humano é possível por uma inteligência artificial? As máquinas poderiam simular, por exemplo, a consciência?

(NC02) Então, a - consciência é algo que é dado. O cérebro ele é holográfico, né? A gente achava que, por exemplo, cada cada estrutura do cérebro tem um especificamente uma área, tem uma certa especificidade, mas o cérebro trabalha com circuitos. É como se fosse entrar num salão, uma um botãozinho acende um monte de de lâmpada, não é uma lâmpada só. Então vai depender muito da educação, quantos circuitos ele conseguiu ir fazendo desde a mãe colocar aquele chocalhozinho lá no no bercinho e depois todos os estímulos que aquela pessoa teve, as experiências de vida, a a a a onde estudou, onde trabalhou, que ambiente percorreu, ela vai fazendo circuitos neuronais, né? Então, são muitas experiências eh eh que o ser humano vai tendo através do tempo que eu acredito que a inteligência ainda não chegou, né? Ela consegue calcular que nem no pro jogo de xadrez, ela calcula muito bem e e ela eu eu passei de um assunto pro outro, não tem problema não, né?

(LPA) Não, não, não.

(NC02) E ela – e ela ajuda muito porque você joga xadrez, né? Então nós estávamos muito condicionados aqueles livros do Ninzovit, do Palau, então dominar o centro, rocar o mais rápido possível, eh eh eh eh coluna aberta, torre na sétima, né?

(LPA) Sim.

(NC02) A inteligência vem mostrar que não precisa ficar preso nesses dogmas que eram muitos dogmas, né? Não tô dizendo que estavam errados, mas que nem o Carlsen calcula diferente.

(LPA) Uhum. É. Eh, eu acho que eh sim, eu eu entendo. E aí vem justamente a pergunta, eh, se, por exemplo, um dos próximos passos que poderia ser dado pela inteligência artificial é colocar justamente essas emoções no processo de tomar decisões.

(NC02) As emoções elas vêm de memórias infantis. Que que é isso? Você assistiu Blade Runner?

(LPA) Sim.

(NC02) Você lembra daquela mocinha que ele apaixona?

(LPA) Sim, sim

(NC02) Ela tinha experiências infantis da sobrinha do cara que fez, né? Você lembra disso?

(LPA) sim, lembro.

(NC02) Então, as os outros não tinham emoções porque eles não tinham memória infantil. Aí o cara percebe que as emoções vêm das memórias infantis. Quer dizer, precisaria ter como se fosse um histórico dentro do da cabecinha do da desses grandes supercomputadores, né?

(LPA) Entendi. E e com relação, por exemplo, a metáforas, analogias, eles são processos também de simulação mental exclusivos dos seres humanos.

(NC02) Eh, no no - por enquanto sim, né? Por enquanto sim, eles calculam muito bem algoritmo, né? Agora, sensações, sensações vem da que nem vem, vem muito do nosso sistema límbico, né? Porque na realidade nós temos três cérebros, né? A gente fala de um cérebro só, mas nós temos três, né? Um em cada andar. Tem tem o cérebro dos répteis, depois tem os dos mais os outros anfíbios, né? E depois aqui vem o nosso cérebro, né, com o é toda uma fundamentação do sistema límbico de dopamina, uma série de neurotransmissores, né, que fazem eh gerar essa todas essas emoções. tem a visão e tem uma coisa é muito diferente, porque nós viemos da ciência, da ciência normal, né, nós falamos em significado das coisas, né? Por exemplo, uma mesa é uma mesa, pronto, cadeira é cadeira, pronto, né? No aspecto de emoção, o Lacan, eu sou lacaniano, estudei 30 anos, né, no genoma humano. Aí ele não fala de significado das coisas, ele fala de significante das coisas. Por exemplo, eu falo mesa, você pensa em quê? –

(LPA) numa tábua com quatro pés.

(NC02) Eu penso num local que eu opero.

(LPA) Ah, sim, sim, sim

(NC02) Um outro cara pensa num local que ele come, no outro pensa num local que ele joga. Então a mesma palavra, enquanto eu falo, você tá pensando numa coisa, eu tô pensando em outra coisa. Em tudo isso, então ele fala que o sujeito surge dessa cadeia de significantes. Então é a teoria do caso único. Cada pessoa é única. Você poderia colocar o computador para entender que mesa é aquele negócio de madeira com como você falou. Então para – chegar nesse nível de especificidade, né?

(LPA) Perfeito.

(NC02) Cada computador teria que ter todo um algoritmos infinitos né?

(LPA) É, e tem uma outra questão, eh, que eu até tenho lido num num livro que fala justamente sobre inteligência e fala desses aspectos, por exemplo, da da questão da recompensa, da dopamina, né? Mas ele chegou num capítulo que eu achei bastante interessante, que ele fala dessa evolução do neocórtex e ele fala que foi justamente devido à sociabilização, ou seja, pode a inteligência artificial se tornar social, interagir umas com as outras, como que as neurociências interpretam esse tipo de questão?

(NC02) o o ser humano, ele ele ele se desenvolveu melhor quando ele passou a viver em comunidade, né? Porque a comunidade é o meu é protege ele, né? Inicialmente das intempéries, né? que eles moravam juntos depois dos animais, dos predadores, os as cidades, todas foram desenvolvendo dessa forma. O homem é um animal social. Um tema inteligente gente estudar também é a inteligência coletiva, né, que é um tema muito bonito. Acho que ele a a inteligência artificial dá para caminhar na inteligência coletiva, né? Agora, para desenvolver o neocórtex é tão é um número tão grande de neurônios que fala que é maior do que do que a galáxia toda, né? A quantidade de conexões que tem. Eu não sei se chegaria nisso, né? No na na minha opinião ainda não, né? –

(LPA) E quanto a própria eh as neurociências, elas elas aprendem também com a inteligência artificial? Tem al... tem alguma coisa assim que que algum exemplo?

(NC02) Aprendem, aprendem. Principalmente que nem em medicina mesmo, né? Você o a pessoa de depois você me manda um um zap, você mandou, né? Eu te mando minha própria cirurgia, você vai ver o o o médico ele ele reaprendeu a operar porque ele abria. Hoje em dia ele ele fica como se fosse um joystick, né? Como se fosse a criança jogando videogame. Por isso que na minha opinião, as crianças aprendendo videogame estão se preparando pro futuro. Essa é a minha opinião, né? Meu filho fica no videogame. Pois é, eu acho que tudo ele vai, nós vamos fazendo com isso aqui. Se você vê o médico, ele tá, são oito braços, um corta, um queima. ele teve a a medicina teve que reaprender com a inteligência artificial, né, mesmo a a ressonância funcional, ela vê o cérebro funcionando. São novas informações que vem com a tecnologia. A tecnologia veio para ficar, né? E ela invade a medicina mesmo, né? invade. Eu sou da época que não tinha. Eu fui fazer engenharia biomédica no Rio de Janeiro em 79. Sabe quantos computadores tinham? Um.

Era como se fosse uma locomotiva. Duas locomotiva. Uma, a pessoa subia a escada, ficava na locomotiva, todo mundo olhando assim. Deve ser um gênio, né? Era DOS, né? E atrás tinha um outra locomotiva que era o refrigerador da primeira. Então, hoje em dia no celular tem muito mais informação que nem eu não tinha tomografia. Hoje em dia eu tenho tomografia. Há muito tempo atrás eu tinha que sair da minha casa para ir ver uma tomografia. Hoje eles me mandam no celular. Quer dizer, você vê quanta evolução que teve, né? E nesses tomógrafos de alta performance, você consegue ver eh por densidade, você pega um joystick e vai vendo por densidade que líquido que é, se é um líquido natural, se é água, se é pus, quer – dizer, eh,

(LPA) Será que a inteligência artificial vai chegar no mesmo patamar do ser humano em termos de inteligência?

(NC02) eu acho que não dá para chegar ainda, né, em algumas áreas que nem na minha área não é tão simples a a inteligência artificial saber tratar, por exemplo, de trauma. Aneurisma até hoje em dia faz embolização, né? né, pela perna coloca uma molinha. Mas no Brasil nós ainda com esse com o carro, né, com esses acidentes de trânsito que a gente tem, que eu falo que é o dinossauro dos nossos tempos, né, eh a pessoa chega com aquilo muito quebrado. Então você tem que montar como se fosse um Lego, tem que abrir aí, tem que montar, tem que pôr o osso ali e coagular. É. é muito específico a neurocirurgia, né? – Principalmente –

(LPA) É que eu acho, é que eu acho que o a inteligência artificial ela adquiriu uma aura, né, principalmente com a vitória do Deep Blue sobre Kasparov. Foi um marco o Kasparov.

(NC02) é ele jogou uma Caro-Kan bem mal. É, ele ele ele tem um livro dele da Caro-Kan que eu li ele não fez nada que ele falou no livro. É. –

(LPA) Mais ou menos. É isso, Dr. (NC02), se o senhor quiser ...

(NC02)Eu conversei muito com o Kasparov, você lembra? É, -

(LPA) acrescentar. Ah, é? Ah, sim. Quando ele veio pro Brasil. É verdade. Ah,

(NC02) conversei, a esposa dele queria que eu fosse para a Rússia, é da palestra lá. Como que você não vai? Falei: "Não, não. Daí eu tinha um cérebro lá, falei: "O senhor quer ver?" Não, não quero não. A esposa dele quis ver. Muito solícita, né? a – esposa dele. E

tem umas frases, só não quero ficar tomando seu tempo, mas umas frases muito bonitinhas, né, do lembro, você lembra do Mauro de Souza, que é o GM,

(LPA) sim.

(NC02) joguei com ele, né, e eu fazia uma tese médica de xadrez e desenvolvimento da inteligência, né, e eu perguntei: "O que que é o xadrez?" Ele falou: "Xadrez é combinação bonito". E perguntei para um menininho jogando com o idoso, o que que é o xadrez para você? Ele falou: "Ô tio, ô tio, xadrez é inguar. Inguar, caçar passarinho. Você tem que saber armar. Eu acho maravilhosa essa frase. E é mesmo, você tem que armar, né? Você tem que deixar uma peça que depois você vai fazer, você vai, você vai armando, né? Eu achei a melhor – definição.

(LPA) E isso e isso a inteligência artificial faz muito bem, né?

(NC02) Faz, mas foi um prazer conversar com você. Desculpe se eu me estendi.

(LPA) Não, imagina. Eu até ia perguntar se o senhor queria acrescentar alguma coisa, falar de algum tópico. Uhum.

(NC02) Não, não. Só, só tenho que agradecer. O xadrez me ajudou demais. Me ajuda muito a ficar concentrado, principalmente quando eu vejo um paciente em CTI. Você precisa ficar atento, precisa enxergar e ver, porque senão você fica olhando e não não percebe, sabe?

(LPA) São os benefícios do jogo de xadrez.

(NC02) É, você precisa – desconfiar, falar: "Será que é só isso? Esse paciente será que tem só isso?" Me deixa atento, sabe? me deixa raciocinando. Se não senta numa internet, você só vê se vão prender o Bolsonaro, se não vão prender o Bolsonaro, se vão fazer isso, se vão fazer aquilo. Assunto repetido que a gente não tem tanto alcance de poder mudar, né? No naquele cálculo de de do xadrez, se faz grandes amigos, se estuda várias saídas, você fala: "Como que a pessoa pensou nisso? -

(LPA) Uhum. É um aspecto. social, não é? É o que eu mais defendo na nas aulas é isso.

(NC02) Sim. E e isso, uma peça sozinha não ganha, precisa da da da do do da trabalhar em conjunto, né? As peças precisam trabalhar harmoniosamente, né? É isso aí. Eu também trabalho em secretaria de saúde, - eu falo sempre, apago o fogo o dia inteiro lá.

O contratado briga com o com o concursado, falo, precisa de todos, todos, precisa. A nossa maior virtude são as nossas diferenças, né? Por isso que cada peça tem a sua função, né?

(LPA) Hum. Muito legal essa frase. Eu vou parar de gravar Dr. Só parar aqui. Ué onde que é aqui o botãozinho? Às vezes, a, às vezes a gente se perde aqui com

(NC02) Tá – bem. Tá bem. É um prazer, viu? No que você precisar me manda que eu respondo.

APÊNDICE G – TCLE 1

UNIVERSIDADE FEDERAL DE SÃO CARLOS

CENTRO DE EDUCAÇÃO E CIÊNCIAS HUMANAS

DEPARTAMENTO DE CIÊNCIAS HUMANAS E EDUCAÇÃO

PROGRAMA DE PÓS-GRADUAÇÃO EM CIÊNCIA TECNOLOGIA E
SOCIEDADE

TERMO DE CONSENTIMENTO LIVRE E ESCLARECIDO

(Resolução CNS 510/2016)

O JOGO DE XADREZ COMO LABORATÓRIO DO PROCESSO DE TOMAR DECISÕES NAS PERSPECTIVAS DA INTELIGÊNCIA ARTIFICIAL E DA INTELIGÊNCIA HUMANA

Eu, Léo Pasqualini de Andrade, estudante do Programa de Pós-Graduação em Ciência Tecnologia e Sociedade da Universidade Federal de São Carlos – UFSCar o (a) convido a participar da pesquisa “O Jogo de Xadrez como Laboratório do Processo de Tomar Decisões nas Perspectivas da Inteligência Artificial e da Inteligência Humana” orientado pela Profa. Dra. Sylvia Iasulaitis e coorientado pelo professor Dr, Alan Demétrius Baria Valejo.

Temos como objetivo conhecer as diferenças da tomada de decisão nas perspectivas de humanos e da Inteligência Artificial no jogo de xadrez. Enquanto a Inteligência Artificial usa algoritmos construídos por cientistas da computação para tomar decisões, a Inteligência Humana necessita de diversos fatores cognitivos, como por exemplo, planejamento, flexibilidade mental, emoções, aprendizados (reconhecimento de padrões), além das abstrações e do puro cálculo de variantes. O jogo de xadrez é um modelo experimental onde podem ser estudados os processos cognitivos humanos, como memória, percepção e resolução de problemas e que também foi utilizado pelos cientistas da computação para desenvolver e construir a inteligência artificial especialista para o jogo. A proposta deste estudo é conhecer e analisar os fatores cognitivos e sociais, ou seja, a construção social que influencia as tomadas de decisão dos seres humanos nas partidas de xadrez e as diferenças com relação a Inteligência Artificial. A coleta de dados é através de entrevistas com neurocientistas, cientistas da computação e mestres de xadrez. Um formulário online (google forms) coletará dados sobre opiniões dos enxadristas em geral sobre as escolhas de lances e de percepção em diagramas de xadrez.

Você está sendo convidado a responder uma entrevista semiestruturada com tópicos sobre diversos aspectos que envolvem conhecimentos de ciência da computação/neurociências/xadrez.

Em nenhum momento será divulgado o seu nome em qualquer fase do estudo e os pesquisadores tomarão todos os cuidados ao seu alcance para preservar sua identidade.

As perguntas não serão invasivas à intimidade dos participantes, entretanto, esclareço que a participação na pesquisa pode gerar desconforto como resultado da exposição de opiniões pessoais em responder às perguntas. Em caso de encerramento da entrevista por qualquer fator, esta será desconsiderada.

Sua participação nesta pesquisa auxiliará na obtenção de conhecimentos que poderão ser utilizados para fins científicos, proporcionando maiores informações e discussões que poderão trazer benefícios para a área da Ciência, Tecnologia e Sociedade.

Os dados coletados serão utilizados apenas nesta pesquisa e os resultados divulgados em eventos e/ou revistas científicas. Na página do nosso grupo de pesquisa, “INTERFACES” na internet, <https://www.interfaces.ufscar.br/>, Os resultados deste estudo poderão ser acessados e acompanhados livremente pelos participantes.

Sua participação nesta pesquisa consistirá em responder a uma entrevista que será gravada, visto que isto possibilita a preservação dos dados e investigação sobre os temas tratados. Os arquivos serão guardados por 5 anos e após este período serão descartados. O tempo estimado de resposta à entrevista é de 20 a 30 minutos e será aplicada pelo pesquisador responsável através de videoconferência, o que permite a escolha do melhor horário para ambos, pesquisador e entrevistado. Não haverá compensação financeira ou custos pela sua participação. De todo modo, será garantida a indenização referente a quaisquer danos oriundos de sua participação na pesquisa, conforme legislação vigente.

Solicito sua autorização para gravação em áudio e vídeo das entrevistas. As gravações realizadas durante a entrevista semiestruturada serão transcritas pelo pesquisador, garantindo que se mantenha o mais fidedigno possível. Depois de transcrita será apresentada aos participantes para validação das informações.

Os riscos relativos à sua participação na pesquisa podem causar desconfortos, sejam sociais, intelectuais ou culturais, como também cansaço mental, ao responder as questões da entrevista. Caso sinta algum desconforto ou cansaço decorrentes da entrevista, você poderá suspender seu consentimento e participação nesta pesquisa.

Os benefícios de sua participação são o de aumentar os conhecimentos científicos para as áreas da ciência da computação, neurociências e xadrez.

A qualquer momento o (a) senhor (a) pode desistir de participar e retirar seu consentimento. Sua recusa ou desistência não lhe trará nenhum prejuízo profissional, seja em sua relação ao pesquisador, à Instituição em que trabalha ou à Universidade Federal de São Carlos. Todas as informações obtidas através da pesquisa serão confidenciais, sendo assegurado o sigilo sobre sua participação em todas as etapas do estudo. Caso haja

menção a nomes, a eles serão atribuídas letras, com garantia de anonimato nos resultados e publicações, impossibilitando sua identificação.

Este projeto de pesquisa foi aprovado por um Comitê de Ética em Pesquisa (CEP) que é um órgão que protege o bem-estar dos participantes de pesquisas. O CEP é responsável pela avaliação e acompanhamento dos aspectos éticos de todas as pesquisas envolvendo seres humanos, visando garantir a dignidade, os direitos, a segurança e o bem-estar dos participantes de pesquisas. Caso você tenha dúvidas e/ou perguntas sobre seus direitos como participante deste estudo, entre em contato com o **Comitê de Ética em Pesquisa em Seres Humanos (CEP)** da UFSCar que está vinculado à Pró-Reitoria de Pesquisa da universidade, localizado no prédio da reitoria (área sul do campus São Carlos). Endereço: Rodovia Washington Luís km 235 - CEP: 13.565-905 - São Carlos-SP. Telefone: (16) 3351-9685. E-mail: cephumanos@ufscar.br. Horário de atendimento: das 08:30 às 11:30.

O CEP está vinculado à **Comissão Nacional de Ética em Pesquisa (CONEP)** do Conselho Nacional de Saúde (CNS), e o seu funcionamento e atuação são regidos pelas normativas do CNS/Conep. A CONEP tem a função de implementar as normas e diretrizes regulamentadoras de pesquisas envolvendo seres humanos, aprovadas pelo CNS, também atuando conjuntamente com uma rede de Comitês de Ética em Pesquisa (CEP) organizados nas instituições onde as pesquisas se realizam. Endereço: SI Página 2 de 3
Via W 5 Norte, lote D - Edifício PO 700, 3º andar - Asa Norte - CEP: 70719-040 - Brasília-DF. Telefone: (61) 3315-5877 E-mail: conep@saude.gov.br.

Dados para contato (24 horas por dia e sete dias por semana):

Pesquisador Responsável: Léo Pasqualini de Andrade

Endereço: Rua Condessa de São Joaquim n. 239 ap 34 São Paulo-SP

Contato telefônico: 11 94059-6886

E-mail: leopasq10@yahoo.com.br

Declaro que entendi os objetivos, riscos e benefícios de minha participação na pesquisa e concordo em participar

Local e data:

Nome do Pesquisador

Nome do Participante

APÊNDICE H – TCLE 2

UNIVERSIDADE FEDERAL DE SÃO CARLOS
CENTRO DE EDUCAÇÃO E CIÊNCIAS HUMANAS
DEPARTAMENTO DE CIÊNCIAS HUMANAS E EDUCAÇÃO
PROGRAMA DE PÓS-GRADUAÇÃO EM CIÊNCIA TECNOLOGIA E
SOCIEDADE

TERMO DE CONSENTIMENTO LIVRE E ESCLARECIDO

(Resolução CNS 510/2016)

O JOGO DE XADREZ COM OLABORATÓRIO DO PROCESSO DE TOMAR DECISÕES NAS PERSPECTIVAS DA INTELIGÊNCIA ARTIFICIAL E DA INTELIGÊNCIA HUMANA

Eu, Léo Pasqualini de Andrade, estudante do Programa de Pós-Graduação em Ciência Tecnologia e Sociedade da Universidade Federal de São Carlos – UFSCar o (a) convido a participar da pesquisa “O Jogo de Xadrez como Laboratório do Processo de Tomar Decisões nas Perspectivas da Inteligência Artificial e da Inteligência Humana” orientado pela Profa. Dra. Sylvia Iasulaitis e coorientado pelo professor Dr, Alan Demétrius Baria Valejo.

Temos como objetivo conhecer as diferenças da tomada de decisão entre humanos e a Inteligência Artificial no jogo de xadrez. Enquanto a Inteligência Artificial usa algoritmos construídos por cientistas da computação para tomar decisões, a Inteligência Humana necessita de diversos fatores cognitivos, como por exemplo, planejamento, flexibilidade mental, emoções, aprendizados (reconhecimento de padrões), além das abstrações e do puro cálculo de variantes. O jogo de xadrez é um modelo experimental onde podem ser estudados os processos cognitivos humanos, como memória, percepção e resolução de problemas e que também foi utilizado pelos cientistas da computação para desenvolver e construir a inteligência artificial especialista para o jogo. A proposta deste estudo é conhecer e analisar os fatores cognitivos e sociais, ou seja, a construção social que influencia as tomadas de decisão dos seres humanos nas partidas de xadrez e as diferenças com relação a Inteligência Artificial. A coleta de dados é através de entrevistas com neurocientistas, cientistas da computação e mestres de xadrez. Um formulário online (google forms) coletará dados sobre opiniões dos enxadristas em geral sobre as escolhas de lances e de percepção em diagramas de xadrez.

Você está sendo convidado a participar de um questionário sobre o jogo de xadrez.

Em nenhum momento será divulgado o seu nome em qualquer fase do estudo e os pesquisadores tomarão todos os cuidados ao seu alcance para preservar sua identidade.

O formulário de perguntas é individual e realizado de forma online. As perguntas não serão invasivas à intimidade dos participantes, entretanto, esclareço que a participação na pesquisa pode gerar desconforto como resultado da exposição de opiniões pessoais em responder às perguntas. Em caso de encerramento do formulário sem o envio por qualquer motivo ou fator, este será desconsiderado.

Sua participação nessa pesquisa auxiliará na obtenção de conhecimentos que poderão ser utilizados para fins científicos, proporcionando maiores informações e discussões que poderão trazer benefícios para a área da Ciência, Tecnologia e Sociedade.

Os dados coletados serão utilizados apenas nesta pesquisa e os resultados divulgados em eventos e/ou revistas científicas. Na página do nosso grupo de pesquisa, “INTERFACES” na internet, <https://www.interfaces.ufscar.br/>, Os resultados deste estudo poderão ser acessados e acompanhados livremente pelos participantes.

Sua participação nesta pesquisa consistirá em responder a um questionário com diagramas de xadrez. Os arquivos serão guardados por 5 anos e após este período serão descartados. O tempo estimado de resposta ao questionário é de aproximadamente 10 minutos. Não haverá compensação financeira ou custos pela sua participação. De todo modo, será garantida a indenização referente a quaisquer danos oriundos de sua participação na pesquisa, conforme legislação vigente.

Os riscos relativos à sua participação na pesquisa podem causar desconfortos, sejam sociais, intelectuais ou culturais, como também cansaço mental, ao responder o questionário. Caso sinta algum desconforto ou cansaço decorrentes do questionário contido no formulário, você poderá suspender seu consentimento e participação nesta pesquisa.

Os benefícios de sua participação são o de aumentar os conhecimentos científicos para as áreas da ciência da computação, neurociências e xadrez.

A qualquer momento o (a) senhor (a) pode desistir de participar e retirar seu consentimento. Sua recusa ou desistência não lhe trará nenhum prejuízo profissional, seja em sua relação ao pesquisador, à Instituição em que trabalha ou à Universidade Federal de São Carlos. Todas as informações obtidas através da pesquisa serão confidenciais, sendo assegurado o sigilo sobre sua participação em todas as etapas do estudo.

Este projeto de pesquisa foi aprovado por um Comitê de Ética em Pesquisa (CEP) que é um órgão que protege o bem-estar dos participantes de pesquisas. O CEP é responsável pela avaliação e acompanhamento dos aspectos éticos de todas as pesquisas envolvendo seres humanos, visando garantir a dignidade, os direitos, a segurança e o bem-estar dos participantes de pesquisas. Caso você tenha dúvidas e/ou perguntas sobre seus direitos como participante deste estudo, entre em contato com o **Comitê de Ética em Pesquisa em Seres Humanos (CEP)** da UFSCar que está vinculado à Pró-Reitoria de

Pesquisa da universidade, localizado no prédio da reitoria (área sul do campus São Carlos). Endereço: Rodovia Washington Luís km 235 - CEP: 13.565-905 - São Carlos-SP. Telefone: (16) 3351-9685. E-mail: cephumanos@ufscar.br. Horário de atendimento: das 08:30 às 11:30.

O CEP está vinculado à **Comissão Nacional de Ética em Pesquisa (CONEP)** do Conselho Nacional de Saúde (CNS), e o seu funcionamento e atuação são regidos pelas normativas do CNS/Conep. A CONEP tem a função de implementar as normas e diretrizes regulamentadoras de pesquisas envolvendo seres humanos, aprovadas pelo CNS, também atuando conjuntamente com uma rede de Comitês de Ética em Pesquisa (CEP) organizados nas instituições onde as pesquisas se realizam. Endereço: SRTV 701, Via W 5 Norte, lote D - Edifício PO 700, 3º andar - Asa Norte - CEP: 70719-040 - Brasília-DF. Telefone: (61) 3315-5877 E-mail: conep@saude.gov.br.

Dados para contato (24 horas por dia e sete dias por semana):

Pesquisador Responsável: Léo Pasqualini de Andrade

Endereço: Rua Condessa de São Joaquim n. 239 ap 34 São Paulo-SP

Contato telefônico: 11 94059-6886

E-mail: leopasq10@yahoo.com.br

Declaro que entendi os objetivos, riscos e benefícios de minha participação na pesquisa e concordo em participar.

Local e data:

Nome do Pesquisador

Nome do Participante

APÊNDICE I – PERGUNTAS CC

Perguntas a serem feitas aos cientistas da computação:

1. As Inteligências Artificiais, poderão simular abstrações, analogias e estratégias para resolver problemas?
2. A Inteligência Artificial depende de descobertas das neurociências sobre o cérebro para se desenvolver?
3. Uma das próximas etapas para a Inteligência Artificial é a simulação das emoções humanas?
4. Segundo palavras de Stephen Hawking em 2014, “o desenvolvimento da inteligência artificial completa pode significar o fim da raça humana”. O que o sr(a) acha deste tipo de declarações sobre os perigos da Inteligência Artificial?
5. Como os jogos, como o xadrez, podem servir de laboratório para as Inteligências Artificiais, em testes com abstrações, intuição, imaginação?
6. A Inteligência Artificial consegue reconhecer as emoções de humanos? Ou poderá reconhecê-las?
7. O que deve ser mais relevante para o desenvolvimento das Inteligências Artificiais daqui para a frente? Que tipo de pesquisas devem acontecer?

APÊNDICE J – PERGUNTAS NC

Perguntas a serem feitas aos neurocientistas:

1. Quais os fatores que diferenciam máquinas e humanos no processo de aprendizagem?
2. Simular o pensamento humano é possível por uma inteligência artificial?
As máquinas poderiam simular a consciência?
3. Pode a Inteligência Artificial se tornar “social”, interagir umas com as outras? Como as neurociências interpretam este tipo de questão?
4. Um dos próximos passos a ser dado pela Inteligência Artificial (IA), é o de colocar emoções para o processo de tomar decisões?
5. As metáforas e analogias são processos de simulação mental exclusivos dos seres humanos?
6. A Inteligência Artificial, para aprender com o ambiente, necessita de um corpo físico?

APÊNDICE K – PERGUNTAS MX

Perguntas a serem feitas aos mestres de xadrez:

1. Qual sua preferência de jogar xadrez, com humanos ou com Inteligência Artificial (engines)?
2. Implementar as Inteligências Artificiais com estratégias é uma possível melhora ou não se justifica, já que mesmo sem fazer planos ela conquista o objetivo final?
3. No jogo de xadrez, a Inteligência Artificial (engines) tem sido utilizada como ferramenta auxiliar para novas ideias e descobertas que envolvem o jogo, pelos mestres de xadrez. Não seria esta a utilidade das *engines*, auxiliar os serviços humanos, ou elas devem ir além?
4. As decisões das *engines* são sempre preferíveis que a dos humanos? Elas são previsíveis? Existe um padrão?
5. Na sua opinião, Kasparov tinha razão em reclamar que houve intervenção humana para as jogadas de *Deep Blue*?
6. O que falta às *engines*, como elas podem melhorar?
7. Quando você está jogando uma partida, como é feita (ou realizada) a sua tomada de decisão? O que você pensa?

APÊNDICE L - SURVEY

Perguntas aos amadores e adeptos do jogo de xadrez:

1. Há quanto tempo você joga xadrez?
 - a. Menos de 1 (um) ano;
 - b. Entre 1 (um) e 5 (cinco) anos;
 - c. Entre 5 (cinco) e 10 (dez) anos;
 - d. Mais de 10 (dez) anos;
2. Você joga torneios de xadrez?
 - a. Sim;
 - b. Não;
3. Já leu algum livro de xadrez?
 - a. Sim;
 - b. Não;
4. Qual seu rating (FIDE, Chess.com ou lichess.org ou outro)?
 - a. Entre 0 e 1.500;
 - b. Entre 1.501 e 1.800;
 - c. Entre 1.801 e 2.000;
 - d. Entre 2.001 e 2.200;
 - e. Mais de 2.201.

Diagramas-problema aos amadores e adeptos do jogo de xadrez:

1. Qual o lance de sua escolha no Diagrama 1?
 - a. Dama branca captura Torre preta em “g8”;
 - b. Bispo branco move para “e5”
 - c. Dama branca captura peão em “d4”;
 - d. Cavalo branco move para “f7”;
 - e. Torre branca captura Torre preta em “g8”;
 - f. Rei branco move para “g3”;
 - g. Rei branco move para “g1”;
 - h. Peão branco move para “g7”;

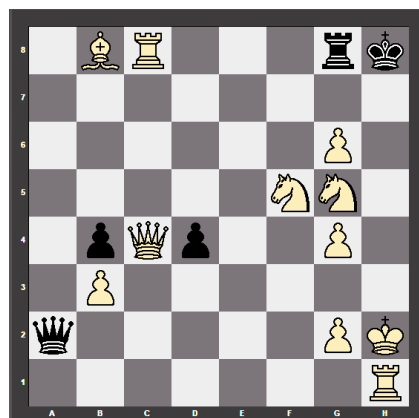


Diagrama 1

2. Qual o lance de sua escolha no Diagrama 2?
- Dama branca move para “e8”;
 - Torre branca move para “e8”;

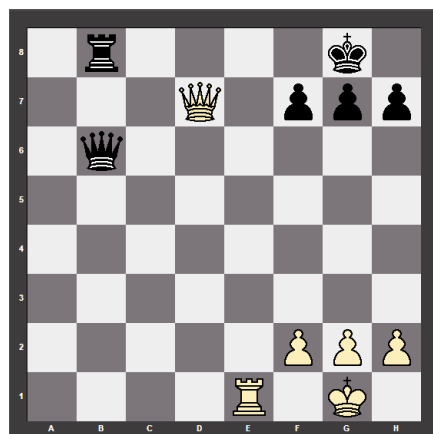


Diagrama 2

3. Qual a primeira ideia que você observa no Diagrama 3?

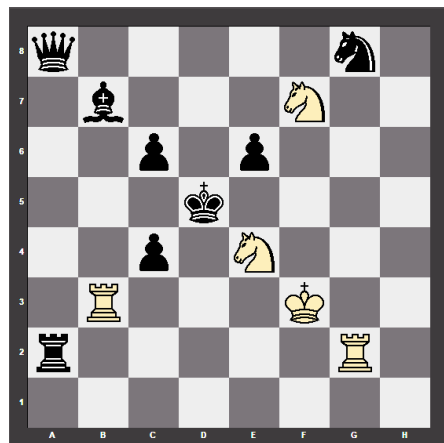


Diagrama 3

4. O Diagrama 4 apresenta um lance (seta vermelha) feito por Magnus Carlsen, ex-campeão mundial (com as peças pretas) contra Ding Liren, atual campeão mundial (que jogava de brancas), em uma partida *online*. Você consegue explicar, em poucas palavras o porquê de Carlsen ter realizado este lance?

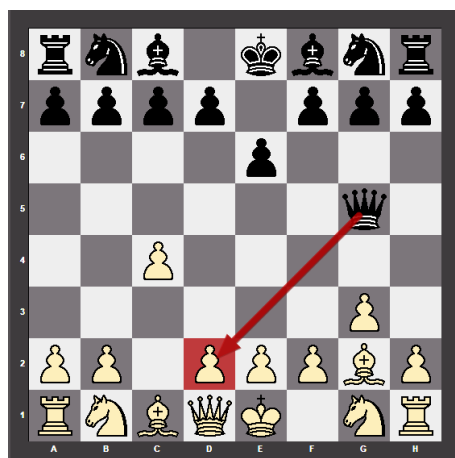


Diagrama 4

5. O Diagrama 5 é da primeira partida pelo campeonato mundial de 1972, entre o russo Boris Spassky e o desafiante, o estadunidense Bobby Fischer. Fischer, com as peças pretas, capturou o Peão branco em “h2” com seu Bispo de “d6”. Os analistas da época comentaram sobre “o incrível erro” cometido por Fischer, que veio a perder a partida, um lance controverso até os nossos dias. Você acha que uma Inteligência Artificial poderia

cometer tal erro? O lance de Fischer foi proposital? O que você pensa a respeito?

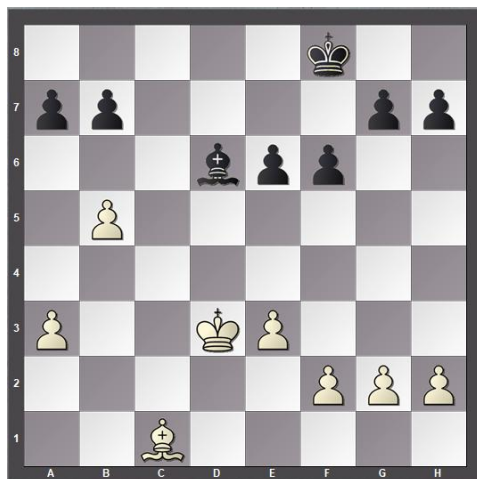


Diagrama 5

6. Ao observar o Diagrama 6, lhe vem alguma lembrança à mente? Qual?

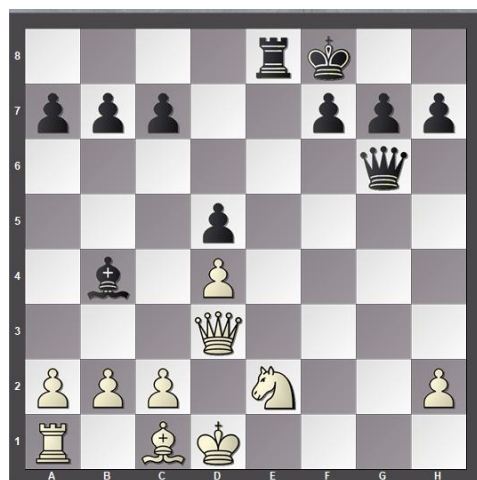


Diagrama 6

7. Faça uma análise do Diagrama 7. Qual o resultado que deve acontecer na partida?
- Vitória das Brancas;
 - Vitória das Pretas;
 - Empate

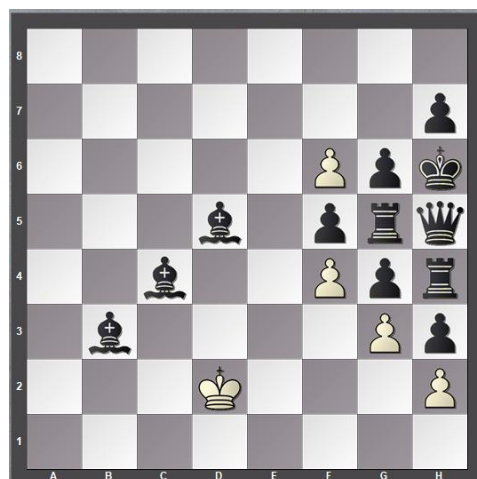


Diagrama 7

8. Faça uma análise do Diagrama 8. Esta posição é de xeque-mate?
- Sim, Dama g1 xeque-mate;
 - Sim, Dama g8 xeque-mate;
 - Sim, Torre d1 xeque-mate;
 - Sim, Dama h7 xeque-mate;
 - Não existe xeque-mate;

```

8 _ _ _ _ _ X _
7 _ _ _ _ Q Q Q
6 _ _ _ _ _
5 _ _ _ _ _
4 _ _ _ _ _
3 _ _ _ _ _
2 _ _ _ _ Z Z Z
1 _ _ W _ E _
  a b c d e f g h

```

Diagrama 8

APÊNDICE M – CARTA CONVITE

Carta Convite de Entrevista

Eu, Professora dra. Sylvia Iasulaitis, docente do programa de pós-graduação em Ciência Tecnologia e Sociedade (PPGCTS), da Universidade Federal de São Carlos (UFSCar), por meio desta, venho solicitar o vosso auxílio para coleta de dados referente a tese de doutorado do aluno sob minha orientação, Léo Pasqualini de Andrade, intitulado “O jogo de xadrez como laboratório do processo de tomar decisões nas perspectivas da inteligência artificial e da inteligência humana”. O objetivo deste estudo é conhecer e analisar as diferenças das tomadas de decisão da inteligência artificial e da inteligência humana no jogo de xadrez.

Os participantes deste estudo compreendem neurocientistas, cientistas da computação e profissionais de xadrez, que serão convidados a responder entrevista semiestruturada e o contato inicial será feito por meio do envio de informações referentes do estudo através de endereço eletrônico, de forma individual e sigilosa, intencionando-se obter uma prévia sobre a sua disponibilidade de participação. O agendamento dos encontros online necessários para a realização de entrevistas semiestruturadas obedecerá às sugestões de datas e horários propostos pelo participante, assegurando que a entrevista não interfira em suas atividades diárias.

O roteiro da entrevista aborda questões sobre a inteligência artificial e as neurociências em suas perspectivas das tomadas de decisão, que intentam subsidiar o alcance geral do estudo. A duração média para a contemplação da entrevista está prevista para um intervalo de 20 a 30 minutos. A identidade de cada um dos sujeitos será preservada. A entrevista será registrada através de videoconferência pelo Google Meet, uma vez obtido o aceite do entrevistado através do Termo de Consentimento Livre e Espontâneo (TCLE).

A pesquisa conta com a aprovação do comitê de ética em pesquisa (CEP) que é um órgão que protege o bem-estar dos participantes de pesquisas. O CEP é responsável pela avaliação e acompanhamento dos aspectos éticos de todas as pesquisas envolvendo seres humanos, visando garantir a dignidade, os direitos, a segurança e o bem-estar dos participantes de pesquisas. Caso você tenha dúvidas e/ou perguntas sobre seus direitos como participante deste estudo, entre em contato com o Comitê de Ética em Pesquisa em Seres Humanos (CEP) da UFSCar que está vinculado à Pró-Reitoria de Pesquisa da universidade, localizado no prédio da reitoria (área sul do campus São Carlos). Endereço: Rodovia Washington Luís km 235 - CEP: 13.565-905 - São Carlos-SP. Telefone: (16) 3351-9685. E-mail: cephumanos@ufscar.br. Horário de atendimento: das 08:30 às 11:30.

O CEP está vinculado à Comissão Nacional de Ética em Pesquisa (CONEP) do Conselho Nacional de Saúde (CNS), e o seu funcionamento e atuação são regidos pelas normativas do CNS/Conep. A CONEP tem a função de implementar as normas e diretrizes regulamentadoras de pesquisas envolvendo seres humanos, aprovadas pelo CNS, também atuando conjuntamente com uma rede de Comitês de Ética em Pesquisa (CEP) organizados nas instituições onde as pesquisas se realizam. Endereço: SRTV 701, Via W 5 Norte, lote D - Edifício PO 700, 3º andar - Asa Norte - CEP: 70719-040 - Brasília-DF. Telefone: (61) 3315-5877 E-mail: conep@saude.gov.br.

Como forma de agradecimento e contrapartida ao vosso auxílio, nos comprometemos a inserir menções de agradecimento ao senhor(a) pelo apoio prestado, tanto na versão final da tese de doutorado quanto em demais outras formas de publicação oriundas do estudo. Certos de vossa compreensão e contribuição para a consolidação da produção científica brasileira, dentro das áreas da ciência da computação, neurociências e para o jogo de xadrez com objeto empírico, agradecemos antecipadamente a atenção disponibilizada.

Data e Assinaturas

Pesquisador Responsável: Léo Pasqualini de Andrade

Contato telefônico: 11 94059-6886

E-mail: leopasq10@yahoo.com.br

Orientadora Sylvia Iasulaitis

Unidade de origem: Ciências Sociais e Ciência de Dados

E-mail: si@ufscar.br

Universidade Federal de São Carlos

Programa de Pós-graduação em Ciência, Tecnologia e Sociedade - PPGCTS

Rod. Washington Luis, km 235 - São Carlos - SP - BR

CEP: 13565-905

Telefones: (16) 3351-8417

E-mail: ppgcts@ufscar.br

ANEXOS

ANEXO A

**PARECER CONSUBSTANCIADO DO CEP – COMITÊ DE ÉTICA EM PESQUISA
COM SERES HUMANOS DA UFSCAR**



PARECER CONSUBSTANCIADO DO CEP

DADOS DO PROJETO DE PESQUISA

Título da Pesquisa: O xadrez como laboratório das tomadas de decisão nas perspectivas da inteligência artificial e da inteligência humana

Pesquisador: LEO PASQUALINI DE ANDRADE

Área Temática:

Versão: 2

CAAE: 74320123.2.0000.5504

Instituição Proponente: Programa de Pós-Graduação em Ciência, Tecnologia e Sociedade

Patrocinador Principal: Universidade Federal de São Carlos/UFSCar

DADOS DO PARECER

Número do Parecer: 6.866.245

Apresentação do Projeto:

As informações elencadas nos campos "Apresentação do Projeto", "Objetivo da Pesquisa" e "Avaliação dos Riscos e Benefícios" foram extraídas do arquivo Informações Básicas da Pesquisa (PB_INFORMAÇÕES_BÁSICAS_DO_PROJETO_2215203, de 14/02/2024) e/ou do Projeto Detalhado (Projeto_de_Pesquisa_detalhado_Leo_Pasqualini_de_Andrade, de 12/02/2024): RESUMO, HIPÓTESE (se houver), METODOLOGIA, CRITÉRIOS DE INCLUSÃO E EXCLUSÃO.

Resumo:

A pesquisa com a inteligência artificial, desde seu início, utilizou o jogo de xadrez como laboratório de estudos. Tomar decisões no jogo de xadrez pelos seres humanos envolve aspectos cognitivos como planejamento, memória e abstrações. Os cálculos de variantes são realizados com muita acurácia pelos algoritmos de inteligência artificial, mas ainda assim apresentam problemas a serem resolvidos. O objetivo deste projeto é conhecer as perspectivas das tomadas de decisão da inteligência artificial e da inteligência humana no jogo de xadrez e analisar suas diferenças. Como método será realizada uma revisão bibliográfica sobre xadrez, tomadas de decisão e inteligência artificial, com o objetivo de conceituar o tema. Em seguida, serão analisadas partidas entre a inteligência artificial e seres humanos no jogo de xadrez. Serão realizadas entrevistas com neurocientistas, cientistas da computação e profissionais de xadrez para discussões sobre o tema. Adeptos do jogo de xadrez serão questionados sobre

Endereço: WASHINGTON LUIZ KM 235

Bairro: JARDIM GUANABARA

UF: SP

Município: SAO CARLOS

CEP: 13.565-905

Telefone: (16)3351-9685

E-mail: cephumanos@ufscar.br



Continuação do Parecer: 6.866.245

posições para se averiguar as tomadas de decisão de lances em determinadas posições. A hipótese é de que os seres humanos são influenciados por aprendizados de aspectos socioculturais, enquanto os algoritmos se baseiam exclusivamente em cálculos, em detrimento de aspectos culturais.

Objetivo da Pesquisa:

As informações elencadas nos campos "Apresentação do Projeto", "Objetivo da Pesquisa" e "Avaliação dos Riscos e Benefícios" foram extraídas do arquivo Informações Básicas da Pesquisa (PB_INFORMAÇÕES_BÁSICAS_DO_PROJETO_2215203, de 14/02/2024) e/ou do Projeto Detalhado (Projeto_de_Pesquisa_detalhado_Leo_Pasqualini_de_Andrade, de 12/02/2024): RESUMO, HIPÓTESE (se houver), METODOLOGIA, CRITÉRIOS DE INCLUSÃO E EXCLUSÃO.

Hipótese:

Enquanto a IA se utiliza de cálculos matemáticos, estatísticas e probabilidades para tomar suas decisões, a IH utiliza fatores cognitivos, como a emoção, abstração e reconhecimento de padrões aprendidos.

Objetivo Primário:

Analisar os processos de tomada de decisão da inteligência artificial especialista e dos seres humanos, em uma partida de xadrez.

Objetivo Secundário:

Conhecer os processos de tomada de decisão da inteligência artificial especialista e dos seres humanos, em uma partida de xadrez.

Avaliação dos Riscos e Benefícios:

As informações elencadas nos campos "Apresentação do Projeto", "Objetivo da Pesquisa" e "Avaliação dos Riscos e Benefícios" foram extraídas do arquivo Informações Básicas da Pesquisa (PB_INFORMAÇÕES_BÁSICAS_DO_PROJETO_2215203, de 14/02/2024) e/ou do Projeto Detalhado (Projeto_de_Pesquisa_detalhado_Leo_Pasqualini_de_Andrade, de 12/02/2024): RESUMO, HIPÓTESE (se houver), METODOLOGIA, CRITÉRIOS DE INCLUSÃO E EXCLUSÃO.

Riscos:

OS RISCOS SÃO REFERENTES A POSSÍVEIS DESCONFORTOS DE ORDEM PSÍQUICA OU FÍSICA QUE ENTREVISTAS (ANEXO 1) E QUESTIONÁRIOS (ANEXO 2) PODEM ACARRETAR AOS

Endereço: WASHINGTON LUIZ KM 235

Bairro: JARDIM GUANABARA

UF: SP

Município: SAO CARLOS

CEP: 13.565-905

Telefone: (16)3351-9685

E-mail: cephumanos@ufscar.br



Continuação do Parecer: 6.866.245

PARTICIPANTES DA PESQUISA.

Benefícios:

DENTRE OS BENEFÍCIOS ESTÃO O DE AMPLIAR OS CONHECIMENTOS A RESPEITO DA INTELIGÊNCIA ARTIFICIAL (IA), SEUS PROCESSOS DE TOMAR DECISÕES E COMO A IA PODE ACARREAR EM MUDANÇAS NA SOCIEDADE.

Comentários e Considerações sobre a Pesquisa:

Trata-se de uma pesquisa que deve seguir os preceitos éticos estabelecidos pela Resolução CNS nº 510 de 2016 e suas complementares.

Considerações sobre os Termos de apresentação obrigatória:

Vide campo "Conclusões ou Pendências e Lista de Inadequações"

Recomendações:

Vide campo "Conclusões ou Pendências e Lista de Inadequações"

Conclusões ou Pendências e Lista de Inadequações:

Agradecemos as providências e os cuidados tomados pelos pesquisadores ao apresentarem a 2ª versão do protocolo de pesquisa ao CEP da UFSCar. Trata-se de análise de resposta ao parecer pendente n.6625787 emitido pelo CEP em 26/01/2024.

Seguem abaixo as pendências listadas no parecer anterior do CEP e seu status (atendida, não atendida, parcialmente atendida).

Pendência 1: Número de participantes da pesquisa (ATENDIDA)

O pesquisador esclarece, com relação à projeção do número de participantes: "Realização de entrevistas semiestruturadas com 8 (OITO) neurocientistas, 8 (OITO) cientistas da computação e 8 (OITO) profissionais de xadrez (ANEXO 1) através de videoconferência a ser gravada, pelo Google Meet e armazenada no google drive. Este procedimento da entrevista, a videoconferência visa facilitar o contato com participantes que estejam em lugares distantes do local de residência do entrevistador. 3) Aplicação através de formulário google de diagramas de xadrez aos aficionados do jogo que aceitarem participar da pesquisa, contendo opções tipo múltipla escolha para solucionar problemas simples sobre lances de xadrez que envolvem aspectos enxadrísticos (ANEXO 2).

Endereço: WASHINGTON LUIZ KM 235

Bairro: JARDIM GUANABARA

UF: SP

Município: SAO CARLOS

CEP: 13.565-905

Telefone: (16)3351-9685

E-mail: cephumanos@ufscar.br



Continuação do Parecer: 6.866.245

Pendência 2: Riscos da pesquisa (ATENDIDA)

O pesquisador esclarece, com relação aos riscos da pesquisa aos participantes: "OS RISCOS SÃO REFERENTES A POSSÍVEIS DESCONFORTOS DE ORDEM PSÍQUICA OU FÍSICA QUE ENTREVISTAS (ANEXO 1) E QUESTIONÁRIOS (ANEXO 2) PODEM ACARRETAR AOS PARTICIPANTES DA PESQUISA." Houve também alteração no TCLE: "OS RISCOS RELATIVOS À SUA PARTICIPAÇÃO NA PESQUISA PODEM CAUSAR DESCONFORTOS, SEJAM SOCIAIS, INTELECTUAIS OU CULTURAIS COMO TAMBÉM CANSAÇO MENTAL, AO RESPONDER O QUESTIONÁRIO. CASO SINTA ALGUM DESCONFORTO OU CANSAÇO DECORRENTES DO QUESTIONÁRIO CONTIDO NO FORMULÁRIO, VOCÊ PODERÁ SUSPENDER SEU CONSENTIMENTO E PARTICIPAÇÃO NESTA PESQUISA."

Pendência 3: Acesso aos resultados da pesquisa (ATENDIDA)

O pesquisador esclarece que haverá acesso aos resultados da pesquisa: "OS RESULTADOS SERÃO DIVULGADOS EM ARTIGOS, CONGRESSOS E NA PÁGINA DO NOSSO GRUPO DE PESQUISA, ¿INTERFACES¿ NA INTERNET, <https://www.interfaces.ufscar.br/>"

Pendência 4: Obtenção de dados dos participantes para a entrevista (ATENDIDA)

O pesquisador esclarece, com relação à obtenção de dados dos participantes, que: "OS CONVIDADOS PARA AS ENTREVISTAS DESTE PROJETO SERÃO OS AUTORES RELEVANTES E DE REFERÊNCIA PROVENIENTES DA REVISÃO BIBLIOGRÁFICA DA PESQUISA, CUJAS PUBLICAÇÕES DE ARTIGOS E LIVROS CONTÉM ENDEREÇOS ELETRÔNICOS (E-MAIL OU SÍTIOS DA INTERNET) PÚBLICOS E DISPONIBILIZADOS PARA QUE O CONTATO COM ELES SEJA POSSÍVEL. ATRAVÉS DA CARTA CONVITE (ANEXO 5) O ENTREVISTADO, POR SUA LIVRE VONTADE E ESCOLHA CONSCIENTE, AUTÔNOMA E ESCLARECIDA, PODERÁ ACEITAR OU DECLINAR PARTICIPAR DA PESQUISA. A CARTA CONVITE MANIFESTA O MÁXIMO RESPEITO E CORDIALIDADE COM O ENTREVISTADO (ANEXO 5). COM RESPEITO A LEI GERAL DE PROTEÇÃO DE DADOS (LGPD), NENHUM ENTREVISTADO SERÁ CONVIDADO PARA A ENTREVISTA, SE NÃO HOVER DISPONÍVEL SEU ENDEREÇO ELETRÔNICO (E-MAIL), EM ARTIGOS LIVROS OU SÍTIOS PESSOAIS DA INTERNET. CONSIDERAMOS QUE SERÃO 8 (OITO) ENTREVISTAS COM PESQUISADORES EM NEUROCIÊNCIAS, 8 (OITO) PESQUISADORES DA CIENCIA."

Tais informações constam do TCLE para a entrevista e do TCLE para o questionário online.

Considerações Finais a critério do CEP:

Diante do exposto, o Comitê de ética em pesquisa - CEP, de acordo com as atribuições

Endereço: WASHINGTON LUIZ KM 235

Bairro: JARDIM GUANABARA

UF: SP

Município: SAO CARLOS

Telefone: (16)3351-9685

CEP: 13.565-905

E-mail: cephumanos@ufscar.br



Continuação do Parecer: 6.866.245

definidas na Resolução CNS nº 510 de 2016, manifesta-se por considerar "Aprovado" o projeto. Conforme dispõe o Capítulo VI, Artigo 28, da Resolução Nº 510 de 07 de abril de 2016, a responsabilidade do pesquisador é indelegável e indeclinável e compreende os aspectos éticos e legais, cabendo-lhe, após aprovação deste Comitê de Ética em Pesquisa:

II - conduzir o processo de Consentimento e de Assentimento Livre e Esclarecido;

III - apresentar dados solicitados pelo CEP ou pela CONEP a qualquer momento;

IV - manter os dados da pesquisa em arquivo, físico ou digital, sob sua guarda e responsabilidade, por um período mínimo de 5 (cinco) anos após o término da pesquisa;

V - apresentar no relatório final que o projeto foi desenvolvido conforme delineado, justificando, quando ocorridas, a sua mudança ou interrupção.

Este relatório final deverá ser protocolado via notificação na Plataforma Brasil.

OBSERVAÇÃO: Nos documentos encaminhados por Notificação NÃO DEVE constar alteração no conteúdo do projeto. Caso o projeto tenha sofrido alterações, o pesquisador deverá submeter uma "EMENDA".

Este parecer foi elaborado baseado nos documentos abaixo relacionados:

Tipo Documento	Arquivo	Postagem	Autor	Situação
Informações Básicas do Projeto	PB_INFORMAÇÕES_BÁSICAS_DO_PROJETO_2215203.pdf	14/02/2024 17:26:14		Aceito
TCLE / Termos de Assentimento / Justificativa de Ausência	TCLE_Anexo_3_e_Anexo_4.pdf	12/02/2024 19:04:36	LEO PASQUALINI DE ANDRADE	Aceito
Outros	Carta_Resposta_versao1.pdf	12/02/2024 18:38:57	LEO PASQUALINI DE ANDRADE	Aceito
Projeto Detalhado / Brochura Investigador	Projeto_de_Pesquisa_detalhado_Leo_Pasqualini_de_Andrade.pdf	12/02/2024 18:37:15	LEO PASQUALINI DE ANDRADE	Aceito
Folha de Rosto	Folha_de_rosto_Leo_Pasqualini_de_Andrade_19_09_2023assinado.pdf	19/09/2023 14:12:32	LEO PASQUALINI DE ANDRADE	Aceito

Situação do Parecer:

Aprovado

Necessita apreciação da CONEP:

Endereço: WASHINGTON LUIZ KM 235

Bairro: JARDIM GUANABARA

UF: SP

Município: SÃO CARLOS

CEP: 13.565-905

Telefone: (16)3351-9685

E-mail: cephumanos@ufscar.br



UNIVERSIDADE FEDERAL DE
SÃO CARLOS - UFSCAR



Continuação do Parecer: 6.866.245

Não

SAO CARLOS, 04 de Junho de 2024

Assinado por:
Sonia Regina Zerbetto
(Coordenador(a))

Endereço: WASHINGTON LUIZ KM 235

Bairro: JARDIM GUANABARA

UF: SP

Município: SAO CARLOS

CEP: 13.565-905

Telefone: (16)3351-9685

E-mail: cephumanos@ufscar.br