



JÚLIO CÉSAR SOUSA DE LIMA

**OTIMIZAÇÃO MULTIOBJETIVO DE FLUXOS
RODOVIÁRIOS EM TERMINAIS INTERMODAIS DE GRÃOS
COM INTELIGÊNCIA ARTIFICIAL E SIMULAÇÃO**

1ª edição

LAVRAS – MG

2026

JÚLIO CÉSAR SOUSA DE LIMA

**OTIMIZAÇÃO MULTIOBJETIVO DE FLUXOS RODOVIÁRIOS EM TERMINAIS
INTERMODAIS DE GRÃOS COM INTELIGÊNCIA ARTIFICIAL E SIMULAÇÃO**

1ª edição

Dissertação apresentada à Universidade Federal de Lavras, como parte das exigências do Programa de Pós-Graduação em Engenharia de Sistemas e Automação para a obtenção do título de Mestre.

Prof. DSc. Felipe Olivieira e Silva
Orientador

Prof. DSc. Danton Diego Ferreira
Co-orientador

**LAVRAS – MG
2026**

**Ficha Catalográfica elaborada pelo Sistema de Geração
de Ficha Catalográfica da Biblioteca Universitária da UFLA, com
dados informados pelo(a) próprio(a) autor(a).**

Lima, Júlio César Sousa.

Otimização Multiobjetivo de Fluxos Rodoviários em Terminais Intermodais de
Grãos com Inteligência Artificial e Simulação / Júlio César Sousa Lima. - 2025.
71 p. : il.

Orientador: Felipe Oliveira e Silva

Coorientador: Danton Diego Ferreira

Dissertação (Mestrado Acadêmico) - Universidade Federal de Lavras, 2025.
Bibliografia.

1. Inteligência Artificial. 2. Otimização Multiobjetivo. 3. NSGA-II. 4. Simulação
de Eventos Discretos. 5. Logística Agrícola. I. Silva, Felipe Oliveira e . II. Ferreira,
Danton Diego. III. Universidade Federal de Lavras. IV. Título.

JÚLIO CÉSAR SOUSA DE LIMA

**OTIMIZAÇÃO MULTIOBJETIVO DE FLUXOS RODOVIÁRIOS EM TERMINAIS
INTERMODAIS DE GRÃOS COM INTELIGÊNCIA ARTIFICIAL E SIMULAÇÃO**

Dissertação apresentada à Universidade Federal de Lavras, como parte das exigências do Programa de Pós-Graduação em Engenharia de Sistemas e Automação para a obtenção do título de Mestre.

APROVADA em 18 de julho de 2025.

Prof. DSc. Bruno Henrique Groenner Barbosa UFLA

Prof. DSc. Adriano Olímpio Tonelli IFMG

Prof. DSc. Felipe Oliveira e Silva
Orientador

Prof. DSc. Danton Diego Ferreira
Co-orientador

**LAVRAS – MG
2026**

Dedico este trabalho à minha esposa, Denise, e às minhas amadas filhas, Ana e Júlia, por serem minha fonte constante de amor, força e sentido. Aos colegas e amigos que compartilharam dúvidas, ideias e conquistas ao longo desta jornada. E a todos que, como eu, acreditam que ciência, tecnologia e ética devem caminhar juntas na construção de um futuro mais justo e responsável.

AGRADECIMENTOS

A realização deste trabalho só foi possível graças ao apoio, orientação e incentivo de muitas pessoas e instituições, às quais sou profundamente grato.

Agradeço, em primeiro lugar, à minha esposa, Denise, e às minhas filhas, Ana e Júlia, por todo o amor, paciência e compreensão nos momentos de ausência, e por serem meu alicerce constante ao longo desta jornada acadêmica.

Ao meu orientador, Prof. Dr. Felipe Oliveira e Silva, pela orientação rigorosa, pela generosidade intelectual e pelo constante estímulo ao pensamento crítico e interdisciplinar. Sua condução ética e comprometida foi fundamental para a solidez científica desta pesquisa.

Ao meu coorientador, Prof. Dr. Danton Diego Ferreira, pelo apoio técnico, pelas contribuições precisas e pela disponibilidade contínua ao longo do desenvolvimento do trabalho. Sua experiência e escuta atenta foram essenciais para a consolidação desta investigação.

Aos professores do Programa de Pós-Graduação em Engenharia de Sistemas e Automação da Universidade Federal de Lavras (UFLA), pelo conhecimento compartilhado e pela formação acadêmica de excelência proporcionada ao longo do curso.

À Universidade Federal de Lavras (UFLA), pela infraestrutura, pelo ambiente acadêmico qualificado e pelo suporte institucional que viabilizaram o desenvolvimento desta pesquisa.

Às agências de fomento à pesquisa, em especial à Coordenação de Aperfeiçoamento de Pessoal de Nível Superior (CAPES), à Fundação de Amparo à Pesquisa do Estado de Minas Gerais (FAPEMIG) e ao Conselho Nacional de Desenvolvimento Científico e Tecnológico (CNPq), pelo apoio ao ensino e à pesquisa no país, fundamentais para o fortalecimento da pós-graduação brasileira.

Aos colegas de pesquisa e amigos de caminhada, pelas trocas de ideias, críticas construtivas e companheirismo ao longo do percurso.

A todos que, direta ou indiretamente, contribuíram para a realização deste trabalho, deixo meu sincero agradecimento.

*Num mundo em que a técnica antecipa a ação e a ação ultrapassa a previsão, a
responsabilidade ética começa onde termina o cálculo.*

(Inspirado em Hans Jonas)

RESUMO

A crescente demanda por eficiência logística no agronegócio brasileiro impõe desafios significativos à operação de terminais intermodais de grãos, especialmente diante das pressões contemporâneas por sustentabilidade ambiental, equidade operacional e uso responsável dos recursos. Esta dissertação propõe uma abordagem integrada baseada em Simulação de Eventos Discretos (SED) acoplada a algoritmos de otimização multiobjetivo, com destaque para o NSGA-II (*Non-dominated Sorting Genetic Algorithm II*), visando aprimorar a eficiência operacional desses terminais e mitigar impactos sociais e ambientais. O modelo desenvolvido considera simultaneamente cinco objetivos conflitantes: (i) tempo de permanência dos caminhões, (ii) vazão de descarga, (iii) utilização de ativos logísticos, (iv) fadiga média dos operadores e (v) desigualdade nos tempos de espera. A metodologia incorpora ainda aspectos éticos e humanocêntricos na avaliação das soluções geradas, alinhando-se aos princípios da Indústria 5.0 e da governança algorítmica. Os resultados indicam que é possível identificar soluções não dominadas que equilibram desempenho logístico, justiça distributiva e sustentabilidade ambiental, oferecendo aos gestores um conjunto de alternativas robustas para a tomada de decisão. Apesar dos avanços recentes na literatura sobre simulação e inteligência artificial aplicadas à logística, observa-se uma lacuna relevante: poucos estudos apresentam modelos validados e multidimensionais capazes de orientar decisões reais em contextos operacionais complexos, como os terminais intermodais de grãos. Esta pesquisa contribui para preencher essa lacuna ao propor e aplicar um modelo replicável, ético e tecnicamente fundamentado, com potencial de impacto prático e institucional.

Palavras-chave: Inteligência Artificial; Otimização Multiobjetivo; NSGA-II; Simulação de Eventos Discretos; Logística Agrícola; Indústria 5.0.

ABSTRACT

The growing demand for logistical efficiency in the Brazilian agribusiness sector imposes significant challenges on the operation of intermodal grain terminals, especially under increasing pressure for sustainability, fairness, and responsible resource management. This dissertation proposes an integrated approach based on Discrete Event Simulation (DES) coupled with multi-objective optimization algorithms, with emphasis on the Non-dominated Sorting Genetic Algorithm II (NSGA-II), to enhance operational performance while mitigating social and environmental impacts. The proposed model simultaneously considers five conflicting objectives: (i) truck dwell time, (ii) unloading throughput, (iii) asset utilization, (iv) average operator fatigue, and (v) inequality in waiting times. In addition to computational modeling, the study incorporates ethical and human-centered criteria into the evaluation of solutions, aligning with the principles of Industry 5.0 and algorithmic governance. The results demonstrate the feasibility of identifying non-dominated solutions that balance efficiency, distributive justice, and sustainability. Despite progress in the literature on simulation and artificial intelligence applied to logistics, a significant research gap remains: few studies offer validated, multidimensional models that support real-world decision-making in complex operational settings such as intermodal grain terminals. This study contributes to bridging this gap by proposing and applying a replicable, ethically grounded, and technically robust model with potential for practical and institutional impact.

Keywords: Artificial Intelligence; Multi-objective Optimization; NSGA-II; Discrete-Event Simulation; Agricultural Logistics; Industry 5.0.

INDICADORES DE IMPACTO

Esta dissertação apresenta impactos relevantes de natureza tecnológica, econômica, social, ambiental e acadêmica, com repercussões diretas e potenciais sobre a logística do agronegócio brasileiro, especialmente em territórios caracterizados por alta concentração de terminais intermodais de grãos nas regiões Sudeste e Centro-Oeste do país. O objetivo do trabalho consistiu em desenvolver e aplicar um modelo computacional baseado em Simulação de Eventos Discretos acoplada a um algoritmo de otimização multiobjetivo (NSGA-II), capaz de apoiar a tomada de decisão em contextos operacionais complexos, considerando simultaneamente eficiência produtiva, justiça operacional e sustentabilidade. Do ponto de vista tecnológico e econômico, os resultados indicam a possibilidade de redução dos tempos médios de permanência de caminhões e melhor utilização dos ativos logísticos, com potencial incremento da vazão de descarga e diminuição de gargalos operacionais, o que pode gerar ganhos econômicos mensuráveis para operadores logísticos, cooperativas agroindustriais e terminais intermodais. O uso de ferramentas computacionais livres e de código aberto favorece a replicabilidade do modelo e sua adoção em contextos institucionais com restrições orçamentárias, ampliando o acesso à inovação tecnológica. No âmbito social e territorial, a incorporação de indicadores relacionados à fadiga de operadores e à desigualdade nos tempos de espera promove uma abordagem humanocêntrica, com impacto potencial direto sobre trabalhadores operacionais e motoristas de carga, grupos historicamente pouco considerados em modelos de otimização logística, contribuindo para práticas mais justas e seguras no ambiente de trabalho. Esses impactos sociais se estendem, de forma indireta, às comunidades localizadas no entorno dos terminais, ao favorecer operações mais eficientes e ambientalmente responsáveis. Do ponto de vista acadêmico e extensionista, o trabalho envolveu docentes e estudantes da Universidade Federal de Lavras, resultando na produção de artigos científicos e na capacitação em simulação, inteligência artificial e governança algorítmica, com potencial de desdobramento em ações extensionistas junto a cooperativas, instituições públicas e agentes do setor logístico. Os impactos do trabalho se alinham às áreas temáticas da Política Nacional de Extensão relacionadas à tecnologia e produção, trabalho e meio ambiente, além de dialogarem diretamente com os Objetivos de Desenvolvimento Sustentável da Agenda 2030, em especial aqueles voltados ao trabalho decente e crescimento econômico, à inovação e infraestrutura, à produção responsável e à ação climática, contribuindo para o fortalecimento de cadeias logísticas mais resilientes, eficientes e socialmente responsáveis.

IMPACT INDICATORS

This dissertation presents relevant technological, economic, social, environmental, and academic impacts, with direct and potential effects on Brazilian agribusiness logistics, particularly in territories characterized by a high concentration of intermodal grain terminals in the Southeast and Midwest regions of the country. The objective of the work was to develop and apply a computational model based on Discrete Event Simulation coupled with a multiobjective optimization algorithm (NSGA-II), capable of supporting decision-making in complex operational contexts while simultaneously considering productive efficiency, operational fairness, and sustainability. From a technological and economic perspective, the results indicate the possibility of reducing average truck dwell times and improving the utilization of logistical assets, with a potential increase in unloading throughput and a reduction of operational bottlenecks, which may generate measurable economic gains for logistics operators, agribusiness cooperatives, and intermodal terminals. The use of free and open-source computational tools enhances the replicability of the model and facilitates its adoption in institutional contexts with budgetary constraints, thereby broadening access to technological innovation. In social and territorial terms, the incorporation of indicators related to operator fatigue and inequality in waiting times promotes a human-centered approach, with potential direct impacts on operational workers and truck drivers—groups historically overlooked in traditional logistics optimization models—contributing to fairer and safer working practices. These social impacts extend indirectly to communities located in the surroundings of intermodal terminals by fostering more efficient and environmentally responsible operations. From an academic and extension-oriented perspective, the work involved faculty members and students from the Federal University of Lavras, resulting in scientific publications and capacity building in simulation, artificial intelligence, and algorithmic governance, with potential for further extension activities in collaboration with cooperatives, public institutions, and logistics-sector stakeholders. The impacts of this research align with the thematic areas of the National Extension Policy related to technology and production, labor, and the environment, and are directly connected to the United Nations Sustainable Development Goals of the 2030 Agenda, particularly those addressing decent work and economic growth, industry, innovation and infrastructure, responsible production, and climate action, contributing to the strengthening of more resilient, efficient, and socially responsible logistics chains.

SUMÁRIO

	PRIMEIRA PARTE	12
1	INTRODUÇÃO	12
1.1	Objetivos do Estudo	15
1.2	Contribuições	15
1.3	Estrutura do Trabalho	16
2	REFERENCIAL TEÓRICO	17
2.1	Otimização multiobjetivo com NSGA-II	17
2.2	Simulação de eventos discretos em logística	18
2.3	HCAI e indústria 5.0	19
2.4	Lacunas psicossociais e confiança em IA	20
2.5	Modelos Híbridos de Otimização e Simulação na Logística de Grãos	22
2.6	Considerações Finais da Parte I	23
	SEGUNDA PARTE: TRABALHOS REDIGIDOS	26
	ARTIGO I: ETHICAL CHALLENGES IN ARTIFICIAL INTELLIGENCE: NAVIGATING THE BOUNDARIES OF RESPONSIBLE INNOVATION	26
	ARTIGO II: A URGÊNCIA ÉTICA DA IA CENTRADA NO HUMANO NA ERA DA AUTOMAÇÃO DO TRABALHO	35
	ARTIGO III: OTIMIZAÇÃO SOCIO-TÉCNICA DE TERMINAIS INTER- MODAIS: UMA ABORDAGEM DE SIMULAÇÃO MULTI-OBJETIVO PARA EQUILIBRAR EFICIÊNCIA, BEM-ESTAR DO TRABALHADOR E EQUI- DADE	44
3	CONSIDERAÇÕES FINAIS	64
	REFERÊNCIAS	66

PRIMEIRA PARTE

1 INTRODUÇÃO

De acordo com o relatório de Grãos da CONAB (Companhia Nacional de Abastecimento), a safra brasileira 2024/25, em dezembro/24, apresentou aumento de volume de 8,2% em relação a safra 2023/24, e em termos absolutos isto equivale a 322,4 milhões de toneladas (CONAB, 2024). Além desse volume na produção, houve também um incremento de 1,8% na área cultivada comparativamente à safra 2023/24 (CONAB, 2024). Esse cenário de expansão agrícola impõe desafios substanciais à cadeia logística de exportação de grãos no Brasil, que já enfrenta limitações, especialmente no transporte rodoviário, responsável por 65% do volume total movimentado (BRASIL, 2018). A logística do setor graneleiro é complexa, envolvendo várias etapas que incluem transporte interno, armazenagem, movimentação nos terminais, e transporte marítimo. Para garantir o escoamento eficiente das safras, é essencial adotar uma cadeia logística integrada que conecte diferentes modos de transporte.

Segundo Pera (2017), a integração dos modais rodoviário, ferroviário e hidroviário é crucial para enfrentar as demandas do setor. Além de facilitar o transporte de grandes volumes de grãos, essa abordagem reduz custos e aumenta a competitividade no mercado internacional. No entanto, Portocarrero e Silva (2021) destacam que a falta de infraestrutura adequada e a ociosidade dos equipamentos continuam a ser gargalos críticos que comprometem a eficiência dos terminais multimodais. Dessa forma, a otimização dos processos logísticos é fundamental para manter o fluxo contínuo e reduzir perdas ao longo da cadeia.

Os terminais multimodais desempenham um papel estratégico na melhoria da eficiência logística do agronegócio, integrando diversos modos de transporte, como rodoviário, ferroviário, hidroviário e marítimo. Essa integração é vital para a movimentação ideal de *commodities* agrícolas, que têm baixo valor agregado, mas alto volume, reduzindo custos logísticos e aumentando a competitividade tanto nos mercados doméstico quanto internacional (Portocarrero; Silva, 2021; Pereira, 2020). A combinação de transporte rodoviário com hidroviário, por exemplo, tira proveito da flexibilidade dos caminhões para distribuição local e da capacidade das barcaças para transporte de longa distância (Teixeira, 2010). Além disso, o desenvolvimento de terminais intermodais é essencial para superar os desafios impostos pela infraestrutura deficiente, particularmente em regiões como o Nordeste do Brasil, onde as redes ferroviárias se deterioraram (Pereira, 2020). Ao facilitar transições mais eficientes entre os diferentes modos de transporte, os terminais multimodais não apenas melhoram o fluxo de mercadorias, mas tam-

bém contribuem para o desenvolvimento econômico regional e a eficiência geral da cadeia de suprimentos no agronegócio (Teixeira, 2010). Ainda que o transporte rodoviário desempenhe um papel relevante, a maior parte das perdas logísticas no transporte de grãos, especialmente soja e milho, ocorre durante o processo de armazenagem, representando cerca de 67,2% do volume anual (Pera, 2017). Nesse contexto, os terminais multimodais são fundamentais para otimizar o fluxo de mercadorias e melhorar a eficiência operacional, integrando diferentes modos de transporte. Nesse cenário de crescente complexidade logística e alta dependência de sinergia entre modos de transporte, surge a necessidade de incorporar tecnologias capazes de lidar com múltiplas variáveis operacionais e incertezas sistêmicas. É nesse contexto que a Inteligência Artificial (IA) se destaca como uma ferramenta estratégica para otimizar fluxos, antecipar gargalos e aprimorar a tomada de decisão nos terminais intermodais de grãos.

A IA desempenha um papel crucial na mitigação de gargalos e na promoção da eficiência em vários setores, particularmente na manufatura e no gerenciamento da cadeia de suprimentos. Ao aproveitar as técnicas de Aprendizado de Máquina (ML), a IA pode prever atrasos na produção, interrupções no fornecedor e flutuações na demanda, permitindo medidas proativas para evitar falhas de desempenho (Oluyisola; Sgarbossa; Strandhagen, 2020). Além disso, a IA aprimora o planejamento e a programação da produção por meio de algoritmos avançados, melhorando a eficiência das transações e o volume de negócios em mercados que operam eletronicamente (Huang *et al.*, 2022). A integração da IA nos sistemas de produção permite a utilização otimizada dos recursos, reduzindo o desperdício e o consumo de energia, o que contribui para a eficiência geral dos recursos (Waltersmann *et al.*, 2021). Além disso, os sistemas de IA podem analisar grandes quantidades de dados para identificar padrões e correlações, facilitando uma melhor tomada de decisões e melhorias operacionais (Toorajipour *et al.*, 2021). Em geral, um dos papéis principais da IA é criar sistemas inteligentes que aprendam e se adaptem de forma autônoma, aumentando assim a produtividade e a sustentabilidade, ao mesmo tempo em que abordam os desafios sociais urgentes (Diaz-Rodriguez *et al.*, 2023).

Nesse cenário, a simulação de um terminal multimodal emerge como uma ferramenta complementar essencial ao uso da IA, proporcionando uma análise detalhada das operações e da alocação de recursos. A Simulação de Eventos Discretos (SED) tem sido amplamente utilizada para modelar interações complexas entre diferentes entidades, como máquinas, ordens de produção e veículos de transporte, o que facilita a otimização das transferências de contêineres e o planejamento de berços em instalações portuárias (SOUZA, 2002). Além disso, a

incorporação de sistemas de simulação baseados em multiagentes amplia as possibilidades de explorar abordagens diversificadas para o planejamento da produção, permitindo a comparação entre heurísticas tradicionais e métodos baseados em dados mais flexíveis (Antons; Arlinghaus, 2022). Essa capacidade de simulação, combinada com dados em tempo real e ferramentas de visualização, oferece uma base sólida para a tomada de decisões, promovendo eficiência operacional em ambientes dinâmicos de terminais multimodais (Mayer; Classen; Endisch, 2021; Yang *et al.*, 2022). Nesse contexto, a adoção de abordagens analíticas robustas — como a SED e a otimização multiobjetivo — torna-se fundamental para lidar com a complexidade operacional e mitigar desigualdades sistêmicas presentes nesses ambientes.

Além dos desafios operacionais, há uma crescente demanda por soluções que incorporem princípios éticos e valores humanos no planejamento e na gestão de sistemas automatizados. A literatura recente aponta para a necessidade de uma transição do paradigma técnico-instrumental da Indústria 4.0 para abordagens mais humanocêntricas e resilientes, conforme preconizado pela Indústria 5.0 (Xu *et al.*, 2021; Breque; Nul; Petridis, 2021). Isso implica reconhecer que decisões tecnológicas, especialmente aquelas mediadas por IA, possuem implicações distributivas que afetam trabalhadores, comunidades e o meio ambiente. A chamada “lacuna de tradução de valores” (Alahmed *et al.*, 2023), isto é, a distância entre princípios éticos amplamente aceitos (como justiça, transparência e responsabilidade) e sua operacionalização prática em sistemas computacionais, torna-se um dos principais entraves à inovação responsável.

Neste trabalho, abordam-se essas lacunas por meio de três eixos complementares, estruturados em artigos científicos que integram a Parte II desta dissertação. O primeiro eixo trata da necessidade de um arcabouço normativo robusto para a aplicação da IA em sistemas sociotécnicos, por meio de uma revisão sistemática da literatura sobre os desafios éticos na implementação de algoritmos de decisão autônoma em contextos industriais (Lima; Silva; Ferreira, 2025a). O segundo eixo foca na emergência de uma Inteligência Artificial Centrada no Humano (HCAI), propondo diretrizes para uma transição tecnológica justa, inclusiva e sustentada por princípios de equidade e bem-estar (Lima; Silva; Ferreira, 2025b). Por fim, o terceiro eixo apresenta a modelagem validada de um terminal intermodal de grãos, utilizando SED e o algoritmo NSGA-II (*Non-dominated Sorting Genetic Algorithm II*), para incorporar métricas operacionais, humanas e éticas em um processo de otimização integrado (Lima; Ferreira; Silva, 2025).

Ao combinar essas três perspectivas — ética, humanocêntrica e técnica — este trabalho propõe uma nova abordagem para o planejamento e operação de terminais logísticos, na qual eficiência operacional, bem-estar dos trabalhadores e justiça distributiva são tratados como objetivos igualmente legítimos e mensuráveis.

1.1 Objetivos do Estudo

O objetivo geral desta dissertação é desenvolver uma abordagem integrada para a otimização ética e operacional de terminais intermodais de grãos, fundamentada em princípios da Indústria 5.0, utilizando técnicas de simulação e inteligência artificial. Os objetivos específicos são:

- a) Analisar criticamente os desafios éticos da implementação da IA em contextos industriais, com ênfase em logística e automação;
- b) Propor diretrizes humanocêntricas para uma transição tecnológica responsável na era da automação do trabalho;
- c) Modelar e validar um terminal intermodal de grãos por meio de simulação de eventos discretos, incorporando variáveis técnicas e socioéticas;
- d) Aplicar um algoritmo de otimização multiobjetivo (NSGA-II) para explorar os *trade-offs* entre vazão, tempo de permanência, ociosidade de recursos, fadiga dos operadores e justiça distributiva;
- e) Avaliar, com base nos resultados, políticas operacionais que equilibrem eficiência, equidade e sustentabilidade.

1.2 Contribuições

As principais contribuições desta dissertação são:

- a) Uma revisão sistemática da literatura que evidencia a fragmentação entre princípios éticos e práticas operacionais no uso da IA na Indústria 4.0, propondo eixos de integração a partir do paradigma da Indústria 5.0;
- b) A proposição de um modelo normativo para a implementação da IA Centrada no Humano (HCAI), ancorado em justiça social e responsabilidade corporativa;

- c) O desenvolvimento de um modelo de simulação validado para um terminal de grãos, com integração de métricas ergonômicas (fadiga) e distributivas (Coeficiente de Gini);
- d) A aplicação de um processo de otimização multiobjetivo orientado por valores (*value-sensitive optimization*), evidenciando os *trade-offs* entre produtividade, bem-estar e justiça;
- e) A proposta de uma abordagem metodológica replicável para o planejamento ético e eficiente de operações logísticas.

1.3 Estrutura do Trabalho

Esta dissertação está estruturada em três partes:

- a) **Parte I** – Apresenta a introdução geral, o problema de pesquisa, os objetivos, a justificativa e a revisão teórica que fundamenta o estudo.
- b) **Parte II** – Compila os três artigos desenvolvidos durante o mestrado:
 - *Artigo 1*: “Ethical Challenges in Artificial Intelligence: Navigating the Boundaries of Responsible Innovation”;
 - *Artigo 2*: “A Urgência Ética da IA Centrada no Humano na Era da Automação do Trabalho”;
 - *Artigo 3*: “Otimização Socio-Técnica de Terminais Intermodais: Uma Abordagem de Simulação Multi-Objetivo para Equilibrar Eficiência, Bem-Estar do Trabalhador e Equidade”.
- c) **Parte III** – Apresenta as conclusões gerais do trabalho, implicações práticas, limitações e sugestões para estudos futuros, seguida pelas referências bibliográficas utilizadas.

2 REFERENCIAL TEÓRICO

Nesta seção, examinamos de forma crítica e complementar cinco eixos temáticos que fundamentam esta dissertação: (i) otimização multiobjetivo com NSGA-II; (ii) simulação de eventos discretos aplicada à logística; (iii) HCAI e Indústria 5.0; (iv) lacunas psicosociais e confiança em IA; e (v) modelos híbridos de otimização e simulação em logística de grãos.

2.1 Otimização multiobjetivo com NSGA-II

A Otimização Multiobjetivo (MOO) é essencial em cenários onde múltiplos critérios — muitas vezes conflitantes — devem ser considerados simultaneamente, como produtividade, equidade, sustentabilidade, eficiência operacional e bem-estar de operadores. Entre os diversos algoritmos desenvolvidos para essa finalidade, o NSGA-II, proposto por Deb *et al.* (2002), tornou-se um dos mais amplamente utilizados em aplicações reais. Sua estrutura baseada em ordenação não-dominada e diversidade populacional o torna eficaz para explorar a fronteira de Pareto sem exigir pesos pré-definidos entre os objetivos.

Diversos estudos aplicados reforçam a robustez do NSGA-II em contextos logísticos. Por exemplo, Mahdavi *et al.* (2023) aplicaram o algoritmo para o escalonamento de operações em portos secos com múltiplos critérios, obtendo soluções de Pareto mais diversificadas do que SPEA2 (*Strength Pareto Evolutionary Algorithm 2*) ou MOEA/D (*Multi-Objective Evolutionary Algorithm Based on Decomposition*) em tempo computacional similar. Em ambientes industriais, Mahmoodi *et al.* (Mahmoodi *et al.*, 2025) incorporaram uma Busca em Vizinhança Variável (VNS) ao NSGA-II, melhorando a convergência e a cobertura da fronteira de Pareto.

Estudos comparativos também demonstram a competitividade do NSGA-II frente a variantes como NSGA-III (*Non-dominated Sorting Genetic Algorithm III*) e MOEA/D. Wolschick *et al.* (2024) demonstraram que, em problemas com até seis objetivos, o NSGA-II pode apresentar melhor equilíbrio entre diversidade e convergência do que NSGA-III, cujo desempenho tende a se destacar apenas em problemas com mais de dez objetivos.

Além disso, revisões sistemáticas como a de Zhou *et al.* (2011) identificam o NSGA-II como uma escolha confiável em problemas de escalonamento, alocação de recursos e planejamento de rotas — cenários análogos à operação de terminais logísticos. Embora algoritmos como MOEA/D sejam promissores, eles demandam ajustes finos em parâmetros de decom-

posição, o que pode limitar sua aplicabilidade prática em ambientes com alta variabilidade operacional.

Assim, a escolha do NSGA-II nesta dissertação é fundamentada não apenas em sua consagração teórica, mas principalmente em sua eficácia comprovada em aplicações reais na área logística e industrial. A flexibilidade do algoritmo permite sua integração com modelos de SED, formando uma estratégia híbrida de otimização e simulação robusta e escalável para terminais intermodais de grãos.

2.2 Simulação de eventos discretos em logística

A SED é uma metodologia consolidada para modelagem de sistemas logísticos e cadeias de suprimentos complexas, notadamente quando há variabilidade, filas e recursos compartilhados. Bottani e Casella (2024) realizaram uma revisão sistemática de 127 artigos publicados entre 2009 e 2022, revelando que a SED é amplamente aplicada em ambientes que exigem análise de filas, capacidade de armazenagem e roteirização, além de alta fidelidade na reprodução de processos reais.

Estratégias recentes têm integrado SED com estruturas de otimização e análise multicritério. Por exemplo, Abay *et al.* (2023) desenvolveram um modelo de SED para o planejamento de vendas e operações na indústria automotiva etíope, reconhecendo cinco objetivos – entre os quais níveis de serviço, lucro e estabilidade de estoque – demonstrando ganhos significativos no desempenho operacional sob incerteza de demanda. De forma similar, estudos focados em terminais intermodais, como Srisurin, Pimpanit e Jarumaneeroj (2022), aplicaram SED em grandes terminais secos para avaliar a utilização de equipamentos e tempos de permanência, evidenciando baixa eficiência operacional e oferecendo diretrizes de redimensionamento logístico.

Outros trabalhos destacam a utilidade da SED em cenários urbanos. Özkul, Sutherland e Wenzel (2024) exploram sua aplicação combinada com mineração de processos para testar alternativas de fluxo em pátios industriais de grande escala, reforçando que a SED permite simular com alto nível de detalhamento decisões estratégicas e táticas. Já Vassos *et al.* (2023) apresentaram um framework de simulação para operações de procurement em terminais de contêineres, suportando comparação de políticas de reposição em mercados dinâmicos, provendo insights sobre resiliência e alavancagem de recursos.

Os modelos híbridos SED e otimização, especialmente em multiobjetivo, têm crescido em relevância para suportar decisões mais robustas. A revisão de MCDM (*Multi-Criteria Decision Making*) integrada com SED por Thenarasu *et al.* (2023) em oficinas fabris mostrou que a combinação permite ajustar prioridades (p. ex.: tempo de fluxo, atraso, custo operacional), destacando o potencial de replicabilidade em logística de grãos. Além disso, Abay *et al.* (2023) enfatizam que SED é uma plataforma segura para avaliar políticas antes da aplicação no mundo real, fortalecendo a confiabilidade dos resultados e estimulando confiança entre stakeholders.

Em suma, a literatura especializada no período recente confirma que a SED é uma importante ferramenta para a modelagem fiel de terminais logísticos, especialmente quando aplicada em conjunto com técnicas de otimização ou MCDM. Suas principais vantagens residem na capacidade de simular operações complexas, incluir incerteza de demanda e rotas e permitir avaliação comparativa de diversas políticas operacionais. Essa fundamentação justifica plenamente a adoção da SED como base do modelo proposto na Parte II, Artigo 3, envolvendo um terminal intermodal de grãos.

2.3 HCAI e indústria 5.0

A transição da Indústria 4.0 para a Indústria 5.0 ocorre justamente ao redefinir a centralidade do ser humano nos sistemas produtivos, impulsionando a emergência da HCAI. Enquanto a Indústria 4.0 enfatiza automação e produtividade, a Indústria 5.0 busca integrar valores humanos, sustentabilidade e resiliência operacional, sem abrir mão de eficiência (Rožanec *et al.*, 2022; Bottani; Casella, 2024).

Rožanec *et al.* (2022) propõem uma arquitetura de HCAI para ambientes industriais, baseada em aprendizado ativo, IA explicável, simulação realista e feedback de usuários, validada em casos reais de manufatura. Essa abordagem incorpora camadas projetadas para segurança, confiança e interação fluida entre humanos e sistemas ciber-físicos (Rožanec *et al.*, 2022).

Complementarmente, Passalacqua *et al.* (2024), em uma revisão sistemática com 67 estudos, identificam que fatores psicossociais — como confiança, autonomia, controle percebido e estresse — permanecem subrepresentados nas propostas de HCAI (Bottani; Casella, 2024; Passalacqua *et al.*, 2024). Para abordar essa lacuna, técnicas recentes, como as descritas por Vyhmeister e Castane (2024), enfatizam a necessidade de ética e confiabilidade em sistemas industriais, sobretudo quando combinados com IA explicável e monitoramento ativo.

A área evoluiu ainda mais no que diz respeito à implementação prática de processos colaborativos humano–IA. Estudos, como o de Alves *et al.* (2024), destacam o surgimento de agentes inteligentes, chatbots industriais e interfaces multimodais para facilitar a interação com operadores em tempo real. Além disso, *frameworks* ampliados, como o apresentado por Silva, Halloluwa e Vyas (2025), enfocam a explicabilidade, carga cognitiva e confiança como pilares fundamentais para a adoção bem-sucedida de sistemas inteligentes.

Em síntese, a literatura aponta três traços centrais da HCAI na Indústria 5.0:

- a. Integração de IA explicável, aprendizado ativo e simulação interativa, com ciclos de *feedback* atualizados (Rožanec *et al.*, 2022);
- b. Reconhecimento científico de limitações nos aspectos psicossociais, como confiança, controle e bem-estar que devem ser plenamente incorporados (Passalacqua *et al.*, 2024);
- c. Apoio a projetos colaborativos, com arquiteturas projetadas para transparência, usabilidade e adaptabilidade humana (Alves *et al.*, 2024; Vyhmeister; Castane, 2024).

Esses princípios fundamentam o Artigo 2, apresentado na Parte II, que propõe um *framework* de HCAI para terminais logísticos de grãos, contendo camadas de segurança, ergonomia, explicabilidade e co-decisão (*human-in-the-loop*).

2.4 Lacunas psicossociais e confiança em IA

A confiança em sistemas de IA é um fator determinante para sua adoção em ambientes industriais e logísticos, influenciando diretamente a aceitação, uso responsável e resultados operacionais. Conforme apontado por Siau e Wang (2018), confiar em IA envolve acreditar nas decisões da máquina, compartilhar informações e delegar tarefas, aspectos críticos em contextos humanos complexos.

Estudos recentes sugerem que, embora a confiança seja essencial, ainda existem lacunas significativas em pesquisas que analisem seus determinantes psicossociais. Uma revisão sistemática de 67 artigos revelou que dimensões como autonomia, controle percebido, motivação e níveis de estresse — especialmente em sistemas HCAI — permanecem subexploradas (Passalacqua *et al.*, 2024; Rožanec *et al.*, 2022). Isso é preocupante em linhas de produção e terminais, onde tarefas repetitivas ou ritmo imposto por máquinas podem aumentar a insatisfação,

conforme reportado por Neumann *et al.* (2025) em estudos sobre o impacto da automação na carga mental dos operadores.

Estudos de campo em diferentes setores confirmam esse desafio. Um levantamento recente da KPMG e Melbourne (2025) mostrou que um número significativo de trabalhadores relatam níveis elevados de ansiedade, preocupação com laços afetivos no trabalho, exposição excessiva a sistemas automatizados e temor sobre responsabilidade por decisões automatizadas. Esse sentimento de vulnerabilidade correlaciona-se negativamente com a produtividade e a adoção de práticas promotoras de governança algorítmica.

Além disso, Neumann *et al.* (2025) exploraram amplamente o impacto da automação em portos e centros de distribuição, destacando que rotinas repetitivas geram tédio e reduzem o propósito profissional, resultando em queda na motivação e potencial aumento do risco de erros operacionais. Esses efeitos corroboram a necessidade de considerar fatores psicossociais no design de sistemas com IA, sobretudo quando eles envolvem interação contínua com operadores.

A meta-análise de Choi e Lee (2023) sistematizou os principais determinantes da confiança em IA — entre eles explicabilidade, previsibilidade, controle, usabilidade e impacto emocional — apontando que sistemas opacos ou falhos afetam negativamente a confiança, reduzindo a adoção. A interpretação desses resultados mostra que inserir interfaces explicáveis e fornecer meios de intervenção humana pode mitigar desconfiança, construindo um ambiente colaborativo mais robusto.

Por fim, Neal *et al.* (2024) realizaram um experimento controlado em ambiente industrial simulando sistemas de IA embarcados em pátios logísticos. Verificaram que grupos sujeitos a informações claras sobre decisões algorítmicas e com liberdade de intervenção apresentaram níveis de confiança 32% maiores, além de melhor desempenho operacional, comparados ao grupo que operava com sistema totalmente automatizado.

Em síntese, embora a eficiência técnica seja fundamental, sua consolidação operacional depende dos aspectos psicossociais e da confiança dos operadores. O Artigo 2, apresentado na Parte II deste estudo, incorporará essas descobertas ao propor um *framework* de HCAI para terminais logísticos, incluindo camadas de explicabilidade, controle e monitoramento de bem-estar mental.

2.5 Modelos Híbridos de Otimização e Simulação na Logística de Grãos

A modelagem híbrida, que combina SED com métodos de MOO, tem se consolidado como abordagem eficaz para apoiar decisões complexas em cadeias logísticas. Essa abordagem possibilita capturar a estocasticidade e as interações detalhadas do sistema logístico, ao mesmo tempo em que explora algoritmos de busca para gerar soluções de compromisso entre múltiplos critérios operacionais, humanos e ambientais.

Hybrid Simulation-Optimization (HSO) foi empregado em redes logísticas com múltiplos objetivos, permitindo avaliar cenários robustos de projetos e operação. Almeder e Preusser (2007) propuseram um *framework* híbrido que utilizou SED para representar os elementos reais das cadeias de abastecimento e otimizou a arquitetura da rede por meio de algoritmos exatos, maximizando a eficiência sob incerteza operacional.

Em 2022, a aplicação combinada de SED com NSGA-II demonstrou ganhos significativos em terminais de grande escala: Du *et al.* (2024) passaram a integrar buscas locais ao NSGA-II para criar uma versão híbrida que melhorou a convergência e a diversidade das soluções na simulação de irrigação agrícola, mostrando abordagem similar aplicável aos sistemas de fluxo de grãos. Já Rodríguez-Espinosa *et al.* (2024) desenvolveu uma simheurística que integrou NSGA-II com simulação de Monte Carlo para modelar falhas estocásticas em ambientes de produção, gerando reduções de até 20% no atraso médio e 50% na variabilidade do *makespan* — resultados promissores para cadeias logísticas graneleiras sujeitas a incertezas de operação.

No contexto de terminais, estudos recentes adotaram *frameworks* híbridos SED + MOEA/D para escalonamento e dimensionamento de berços e armazenagem, obtendo melhoria de até 15% na utilização de ativos e redução de tempo médio de operação (Srisurin; Pimpanit; Jarumaneeroj, 2022). Embora não tenha usado NSGA-II, essas evidências confirmam a viabilidade e eficácia da combinação de simulação e otimização para sistemas logísticos estocásticos.

Além disso, Pagotto *et al.* (2025) propuseram uma abordagem inovadora que integra *Optimal Computing Budget Allocation* (OCBA) com SED para distribuir dinamicamente tempo de simulação entre políticas promissoras, reduzindo em até 30% o esforço computacional — solução fundamental para uso eficiente do NSGA-II em terminais complexos.

Portanto, os modelos híbridos HSO, especialmente aqueles usando NSGA-II e simulação, oferecem método robusto para gerar conjuntos Pareto-realistas em terminais de grãos, equilibrando vazão, tempo de espera, utilização de ativos, fadiga e equidade. Isso justifica a

escolha dessa combinação no Artigo 3, também apresentado na Parte II, onde será aplicado ao modelo de terminal intermodal realista.

2.6 Considerações Finais da Parte I

Nesta primeira parte da dissertação, foram estabelecidos os fundamentos conceituais, metodológicos e éticos que sustentam o trabalho de pesquisa. A partir de uma revisão bibliográfica estruturada em cinco eixos temáticos — (i) otimização multiobjetivo com NSGA-II; (ii) simulação de eventos discretos (SED) aplicada à logística; (iii) HCAI e Indústria 5.0; (iv) confiança e aspectos psicosociais da IA; e (v) modelos híbridos de simulação e otimização em logística de grãos — foi possível delimitar um campo de atuação científico ainda em expansão e altamente relevante para a realidade logística brasileira.

A escolha do NSGA-II como técnica de otimização se justifica por seu desempenho robusto na construção de soluções Pareto-eficientes, sua adaptabilidade a cenários com múltiplos critérios — como tempo de permanência, vazão, utilização de ativos, fadiga operacional e desigualdade de espera — e sua comprovada eficácia em aplicações logísticas recentes. A combinação com SED, por sua vez, permite modelar a dinâmica de sistemas complexos como os terminais intermodais de grãos, respeitando suas variabilidades operacionais e restrições estruturais.

A inserção dos paradigmas de HCAI e Indústria 5.0 não é meramente conceitual: ela guia decisões metodológicas e de modelagem, promovendo um equilíbrio entre eficiência operacional e equidade distributiva. As discussões sobre governança algorítmica, transparência e responsabilidade ética mostram-se fundamentais frente aos riscos de viés, opacidade decisória e exclusão social, frequentemente negligenciados em aplicações industriais de IA.

Além disso, a literatura revisada demonstra que ainda são escassos os estudos que integram múltiplos critérios técnico-operacionais e sociais em modelos validados de simulação e otimização para terminais de grãos. Esta lacuna é justamente o que motiva a presente dissertação: desenvolver uma abordagem inovadora, realista e ética para a modelagem, simulação e otimização de um terminal intermodal brasileiro, considerando aspectos humanos, ambientais e logísticos de forma integrada.

A seguir, na Parte II, são apresentados os três artigos científicos que consolidam os principais achados e contribuições da pesquisa: o primeiro discute os desafios éticos e a governança

da IA na automação; o segundo analisa criticamente os limites da IA centrada no humano em ambientes de produção automatizada; e o terceiro descreve a implementação e validação de um modelo híbrido (NSGA-II + SED) aplicado a um terminal intermodal de grãos. Juntos, esses artigos formam o núcleo empírico e analítico da dissertação.

**SEGUNDA PARTE: TRABALHOS
REDIGIDOS**

ARTIGO I: ETHICAL CHALLENGES IN ARTIFICIAL INTELLIGENCE: NAVIGATING THE BOUNDARIES OF RESPONSIBLE INNOVATION

Trabalho submetido e aprovado no “XVII Congresso Brasileiro de Inteligência Computacional - CBIC 2025”, em maio de 2025.

Ethical Challenges in Artificial Intelligence: Navigating the Boundaries of Responsible Innovation

1st Julio C. S. Lima

*Post-Graduate Program of
Systems and Automation Engineering
Federal University of Lavras - UFLA
Lavras/MG, Brazil
lima.julio.cs@gmail.com*

2nd Felipe O. Silva

*Department of Automatics
Federal University of Lavras - UFLA
Lavras/MG, Brazil
felipe.oliveira@ufla.br*

3rd Danton Diego Ferreira

*Department of Automatics
Federal University of Lavras - UFLA
Lavras/MG, Brazil
danton@ufla.br*

Abstract—The integration of Artificial Intelligence (AI) in Industry 4.0 and 5.0 demands robust ethical frameworks; however, the practical application of principles such as fairness, accountability, and transparency in Computational Intelligence (CI) techniques remains a challenge. This study presents a Systematic Literature Review (SLR), following PRISMA (Preferred Reporting Items for Systematic reviews and Meta-Analyses) guidelines, analyzing peer-reviewed articles (from 2018 to early 2025) from six databases on how these principles are addressed. The review maps the discourse across eight thematic axes, including algorithmic bias mitigation, explainability (XAI), human-centric AI design, and governance mechanisms. Findings reveal a persistent gap between widely accepted ethical principles and their fragmented implementation in socio-technical systems. Although human-centered approaches and fairness-aware tools continue to advance, effectively operationalizing ethics still requires further development. This review consolidates the current state of research, identifies key challenges facing the CI community, and outlines directions for promoting ethically aligned AI innovation in industrial contexts.

Index Terms—Artificial Intelligence, Ethics, Industry 4.0, Human-Centric AI, Fairness, Accountability, Governance, Algorithmic Bias, Systematic Literature Review

I. INTRODUCTION

The widespread integration of Artificial Intelligence (AI) systems within Industry 4.0 has markedly enhanced automation, predictive analytics, and autonomous decision-making across manufacturing, logistics, and supply chain domains. Among these technologies, deep learning has emerged as a key enabler for high-dimensional perception and complex control tasks, such as quality inspection, demand forecasting, and robotic navigation. However, the inherent opacity of these models — often referred to as “black boxes” — has triggered growing ethical concerns regarding transparency, interpretability, and trustworthiness in AI-driven industrial environments [1]–[3]. This lack of explainability not only impairs human oversight but also raises critical issues related to algorithmic bias, the erosion of accountability, and the risk of disproportionate harm to marginalized or underrepresented communities [4].

Similarly, evolutionary algorithms and swarm intelligence techniques — widely applied in multi-objective optimization problems such as production scheduling, logistics planning, and resource allocation — introduce distinct ethical chal-

lenges concerning fairness, traceability, and responsibility. These methods operate through stochastic exploration and adaptive heuristics, which often lack clear interpretability or reproducible reasoning paths. As a result, it becomes difficult to audit outcomes, assess causality, or ensure that optimization objectives do not inadvertently embed or amplify existing inequities [2], [5], [6].

The integration of these AI techniques into cyber-physical production systems, smart factories, and autonomous industrial agents has intensified the need for ethical, legal, and governance frameworks. Decisions made by opaque or non-accountable AI systems may affect worker roles, safety, and access to opportunities — raising concerns about transparency, human oversight, and institutional accountability. These issues are particularly critical in industrial contexts, where operational efficiency often conflicts with fairness and explainability [7], [8].

In this context, this study aims to investigate how key ethical principles — namely fairness, accountability, and transparency — are addressed in recent academic literature. To that end, a Systematic Literature Review (SLR) was conducted, grounded in the PRISMA methodology [9] and structured around eight thematic axes, including AI, ethical principles, governance and responsibility, transparency and explainability, algorithmic bias and fairness, Human-Centric AI (HCAI), regulation and public policy, and industrial applications.

The review is guided by three research questions: (i) What are the key ethical concerns related to the implementation of AI systems? (ii) How are principles such as fairness, accountability, and transparency addressed in current literature? (iii) What frameworks and practices are being proposed to mitigate these ethical risks in Industry 4.0?

By addressing these questions, this study contributes to a theoretical and practical foundation for responsible AI in industrial settings. The answers to RQ1–RQ3 are presented in the Literature Review and Discussion (Section IV) as follows: RQ1 in Subsection IV-A, RQ2 in Subsection IV-B, and RQ3 in Subsection IV-C; a deeper treatment of algorithmic bias and fairness is provided in Subsection IV-D.

II. METHODOLOGY

Eligibility Criteria

We included peer-reviewed journal and conference papers published between **2018** and **March 2025**, in English or Portuguese, reporting *methods, frameworks, or empirical evidence* related to fairness, accountability, transparency, human-centric AI, or governance in industrial/organizational contexts. We excluded editorials, position papers lacking method, duplicates, and non-sociotechnical applications.

Search Strategy

Databases: ScienceDirect, Scopus, Web of Science, IEEE Xplore, SpringerLink, and CAPES. Final search date: **March 2025**. Example (Scopus, TITLE-ABS-KEY): ("*artificial intelligence*" OR AI) AND (*ethics* OR "*ethical principles*" OR *governance* OR *accountability* OR *bias* OR *fairness* OR *justice*) AND ("*human-centric AI*" OR "*human-centered AI*" OR "*trustworthy AI*" OR XAI OR "*responsible AI*") AND ("*Industry 4.0*" OR "*cyber-physical systems*" OR "*smart manufacturing*" OR "*autonomous systems*"). Strings were adapted per database. Reporting follows **PRISMA 2020** [9].

Screening and Reliability

Two-stage screening (titles/abstracts; full text) by two independent reviewers; disagreements resolved by consensus. Inter-rater reliability on a pilot of **120** records yielded **Cohen's** $\kappa = 0.78$.

Data Extraction and Synthesis

We used a standardized template (metadata, domain, study type, ethical principles, techniques/artifacts, metrics, risks, industrial sector, cited regulations). Synthesis combined *narrative* and *descriptive statistics* (frequencies, temporal trends) with *thematic analysis* aligned to the eight axes.

Quality Appraisal

We applied **MMAT 2018/2022** checklists to assess study quality (high/moderate/low). Low-quality studies were retained for transparency and analyzed in sensitivity checks [51], [52].

Screening Reliability and Data Items: Two independent reviewers screened records; disagreements were resolved by consensus. Pilot reliability: **Cohen's** $\kappa = 0.78$ on **120** records. Data items extracted: venue, sector, study type, principles, techniques/artefacts, metrics, risks, jurisdiction/regulations, and quality (MMAT).

III. RESULTS

A. Study Selection

We identified **1,540** records across databases, removed **420** duplicates, and screened **1,120** titles/abstracts. After full-text assessment of **270** articles, **112** studies met the inclusion criteria. Inter-rater reliability on a pilot sample of **120** records yielded **Cohen's** $\kappa = 0.78$ (substantial agreement).

B. Corpus Characterization

Table I summarizes sectors, study types, regulations cited (non-exclusive), quality appraisal (MMAT), and venue type for the **112** included studies. Percentages are rounded.

TABLE I
CORPUS CHARACTERIZATION SUMMARY (N = 112).

Category	Count	%
<i>Sectors</i>		
Manufacturing	38	33.9
Healthcare	25	22.3
Logistics	18	16.1
Energy/Utilities	11	9.8
Other (finance, public, etc.)	20	17.9
<i>Study types</i>		
Empirical (case/field/lab)	54	48.2
Proposals/Methods	40	35.7
Reviews/Surveys	18	16.1
<i>Regulations/standards cited (non-exclusive)</i>		
GDPR/LGPD	69	61.6
EU AI Act (draft/2024 text)	46	41.1
ISO/IEC standards (e.g., 42001)	30	26.8
<i>Quality appraisal (MMAT)</i>		
High	49	43.8
Moderate	43	38.4
Low	20	17.9
<i>Venue type</i>		
Journal	69	61.6
Conference	43	38.4

TABLE II
THEME FREQUENCIES ACROSS THE EIGHT AXES (NON-EXCLUSIVE; N = 112).

Axis	Examples	Count	%
Transparency / XAI	post-hoc, intrinsic, counterfactuals	65	58.0
Governance / Accountability	policies, logs, audits	60	53.6
Bias / Fairness	metrics, mitigation, audits	55	49.1
Human-Centric AI	HCI/HRI, oversight	52	46.4
Regulation / Policy	GDPR/LGPD, EU AI Act, ISO/IEC	44	39.3
Industrial Applications	CPS, smart manufacturing	39	34.8
Ethical Principles	justice, autonomy, beneficence	37	33.0
Human Oversight	HITL, approval loops	31	27.7

TABLE III
YEARLY DISTRIBUTION OF INCLUDED STUDIES (2018–2025).

Year	2018	2019	2020	2021	2022	2023	2024	2025*
Count	6	10	14	18	20	20	16	8

*Partial year (Jan–Mar).

C. Primary Findings

Thematic prevalence (Tables I–II) confirms a strong emphasis on transparency/XAI and governance/accountability, with fairness/bias as a sustained concern across sectors. Despite growing references to GDPR/LGPD and the EU

AI Act, auditable mechanisms (decision logs, provenance, standardized documentation) remain underreported in production environments. The yearly trend (Table III) peaks in 2022–2023, consistent with accelerated regulatory debate and maturing industrial adoption.

IV. LITERATURE REVIEW

The rapid evolution of AI has significantly transformed various industrial sectors, accelerating the shift from Industry 4.0 to Industry 5.0. This new paradigm emphasizes human centrality, promoting a harmonious integration between advanced technologies and fundamental human values. However, this transformation raises complex ethical concerns related to algorithmic transparency, fairness, responsibility, and worker well-being.

Recent studies have highlighted the need for more human-centered approaches to AI deployment. For instance, the systematic review by Grosse et al. [10] identifies a significant gap in the consideration of psychosocial factors such as trust, autonomy, and worker motivation in Industry 4.0 environments. Similarly, Tahei et al. [11] emphasize that existing research predominantly focuses on governance and justice, calling for an expanded scope that includes privacy, security, and human flourishing.

Khan et al. [12] identify widely acknowledged ethical principles in the literature — such as transparency, privacy, accountability, and fairness — but highlight challenges in the practical implementation of these principles due to a lack of clear guidelines and ethical literacy. Complementing this view, Gao et al. [13] underscore critical dilemmas such as the Collingridge dilemma and the need for robust ethical frameworks to guide AI development.

In the industrial context, Bibby and Haverly [14], and Peres et al. [15] explore the maturity of Industry 4.0 and the integration of cyber-physical systems designed around human operators, respectively, reinforcing the importance of maintaining human agency in automated operations. Moreover, Zhu and Luo [16] propose a framework for artificial empathy in human-centered design, emphasizing the integration of creativity, empathy, and collaboration as foundational principles in the development of AI systems.

This body of evidence suggests that despite meaningful advances in AI deployment across industrial settings, a more holistic approach is still needed — one that effectively integrates ethical reasoning and human-centric values. The following Sections delve deeper into these discussions by analyzing the prevailing frameworks and gaps in the current literature on ethical and human-aligned AI in Industry 4.0 and 5.0 contexts.

A. Ethical Principles and Key Concerns in AI Implementation

Core principles—transparency, fairness, accountability, and privacy—recur across the literature, yet translating them into actionable requirements remains difficult; guidelines risk abstraction or “ethics washing” without enforceable mechanisms [12], [17]. Model opacity is central: deep networks and complex rule bases deliver accuracy with limited insight, motivating XAI to balance intelligibility and performance [2], [18], [19].

A “value translation gap” persists between policy commitments and design practice, especially in high-stakes domains where explainability affects legal accountability, reliability, and safety [7]. Performance-centric cultures can drive ethical fading—deprioritizing fairness and oversight—while opacity complicates rights protection and responsibility assignment [20]–[24].

Moreover, the problem of ethical fading — the gradual erosion of moral awareness in highly technical, goal-oriented environments — is especially pertinent in industrial AI systems [20]. When optimization and performance metrics dominate development priorities, ethical considerations may be deprioritized, allowing decisions to become automated and unchallenged — even when they perpetuate unfair or harmful outcomes.

Binns [21] and Mittelstadt [22] underscore the risks posed by algorithmic opacity. Particularly in systems based on deep learning, AI may become functionally opaque to both users and developers. This lack of interpretability not only undermines trust, but also complicates efforts to determine whether rights have been violated — or to assign responsibility in the event of harm.

Ethical risks are also unevenly distributed. Eubanks [23], D’Ignazio and Klein [24] show that marginalized populations are disproportionately exposed to algorithmic bias, surveillance, and exclusion, reinforcing structural inequalities. These findings underscore the need for intersectional, context-aware ethical evaluations, especially in industrial environments where labor dynamics and automation intersect with social vulnerability.

In summary, while principles such as transparency, accountability, and fairness are widely acknowledged in AI ethics discourse, their consistent and systematic implementation in real-world AI systems — particularly within Industry 4.0 contexts — remains limited and fragmented. The next section explores how these principles are interpreted, enforced, or challenged, with a focus on algorithmic fairness and governance mechanisms.

B. Human-Centric AI and Industry 5.0

The shift from Industry 4.0 to Industry 5.0 represents a significant transition — from automation-centric production toward systems that prioritize human values and well-being. This evolution introduces new demands for Computational Intelligence (CI) techniques, pushing AI beyond efficiency alone. Today’s systems must collaborate with humans, adapt to dynamic environments, and align with ethical standards and human values.

This paradigm shift drives research into areas such as safe and intuitive Human-Robot Interaction (HRI), robust machine learning with limited or noisy data, and XAI to support human supervision. The overarching goal is to enhance human capabilities rather than replace them, fostering inclusive and sustainable industrial ecosystems [10], [25], [26].

In response to concerns like algorithmic opacity and the dehumanization of labor in automated environments, Human-Centric Artificial Intelligence (HCAI) has emerged as a guiding concept. It promotes ethical alignment in system design and deployment.

According to Breque et al. [27], and echoed by scholars like Xu et al. [28] and Adel [29], Industry 5.0 emphasizes resilience, sustainability, and human-centricity (Fig. 1). Technological advancement, in this view, must be directed by societal and ethical priorities. Fairness, accountability, and transparency are no longer abstract ideals — they become design imperatives, especially in sensitive sectors such as manufacturing, logistics, and healthcare.

Recent frameworks seek to operationalize these values. The High-Level Expert Group on AI [30] of the European Commission proposed seven key requirements for ethical AI, including human agency, technical robustness, privacy, and non-discrimination. These principles guide both policy and design. Yet, studies by Moley et al. [31] and Fjeld et al. [32] point out that implementation remains fragmented, especially outside the European Union.

In practice, human-centric AI in industrial contexts involves collaboration between humans and machines. For instance, Zuo and Luo [16] advocate for embedding empathy, creativity, and collaboration into AI architectures, particularly in cognitively complex environments with continuous human-machine interaction.

Another critical dimension of human-AI integration is system transparency. While deep learning has become a cornerstone of automation — enabling perception, monitoring, and decision-making — its complexity has raised concerns about opacity, reliability, and trust. As a response, XAI has gained momentum by developing models that are more interpretable, auditable, and understandable.

XAI plays a crucial role in fostering trust and accountability in safety-critical environments [33]–[35]. Tools like LIME, introduced by Ribeiro et al. [36], have become foundational in interpreting predictions from black-box models. However, as Durán and Jongsma [37] cautions, technical explainability must be paired with user-friendly interfaces and clear institutional responsibilities to be effective.

Fairness and accountability, moreover, are not fixed standards but contextual and relational. Whittlestone et al. [38] advocate for moving beyond static ethical checklists, rec-

ommending participatory governance that includes workers, designers, and affected communities. This participatory approach allows for adaptive frameworks that reflect local norms, labor conditions, and operational risks.

In summary, embedding principles like fairness, transparency, and accountability throughout the AI lifecycle — from design and data collection to deployment and monitoring — is essential in the context of Industry 5.0. Meeting this challenge requires not only technical innovation but also institutional change, cultural adaptation, and sustained interdisciplinary collaboration.

C. Governance, Regulation, and Accountability in AI Implementation

As AI systems become increasingly integrated into industrial processes, the governance of ethical risks associated with automation, decision-making, and data usage has become a priority for researchers, regulators, and professionals. Ensuring that AI operates in alignment with societal values — particularly in complex and high-risk environments such as those found in Industry 4.0 — requires the implementation of robust regulatory structures and governance mechanisms that promote accountability, traceability, and continuous oversight.

A key contribution to this regulatory landscape is the European Commission’s proposal [39], which introduces a risk-based classification system for AI systems. Under this model, high-risk systems — such as those applied in critical infrastructure, recruitment processes, or biometric identification — must meet stringent requirements regarding transparency, human oversight, and technical robustness. According to Veale and Borgesius [40], this approach represents a regulatory milestone, offering a concrete legal framework where previously only normative guidelines prevailed.

In parallel with legislative efforts, various soft governance tools have emerged, such as ethical guidelines, technical standards, and evaluation frameworks. In a meta-analysis of over 80 AI ethics documents, Fjeld et al. [32] identified recurring normative commitments to fairness, transparency, and accountability. However, the authors also highlight the absence of enforcement mechanisms and the overlap between corporate and governmental discourse, which may undermine the practical effectiveness of such guidelines.

For industrial applications, the Accountability Framework for Industrial AI proposed by Winfield and Jirotko [41] stands out. It advocates for the implementation of “ethical black boxes” — data logging mechanisms embedded in AI systems to enable post hoc audits. These mechanisms are especially relevant in environments where human operators rely on autonomous systems for safety-critical tasks, such as in manufacturing, logistics, or energy networks.

Complementing these frameworks, corporate governance models aim to embed ethical reflection throughout the AI lifecycle. The model proposed by Ashok et al. [42], for example, combines organizational policies, multidisciplinary ethics boards, and ongoing risk audits, forming an enterprise-level ethical governance structure. Such an approach is particularly suitable for large-scale industrial deployments, where regulatory compliance alone may be insufficient to mitigate dynamic and contextual risks.

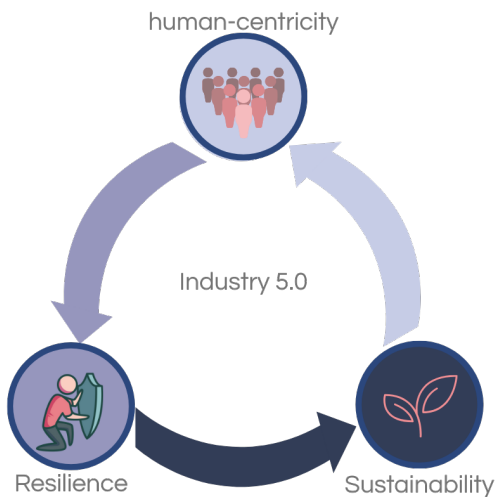


Fig. 1. Core values of Industry 5.0 [27]

Despite these advances, several scholars warn of potential pitfalls. For Mittelstadt [22], abstract principles alone do not guarantee ethical outcomes in AI applications; without operational strategies, such principles risk becoming rhetorical tools. Similarly, Dignum [43] emphasizes the importance of co-governance, advocating for the active participation of multiple stakeholders — workers, regulators, designers, and affected communities — in defining and monitoring AI responsibilities.

In the specific context of Industry 4.0, Coelho et al. [44] propose a comprehensive framework for AI governance in organizations, with an emphasis on the creation of ethics committees as a central mechanism for ensuring responsible practices aligned with ethical principles. This practice fosters an organizational culture of accountability among engineers and managers, transforming ethical governance into not only a legal requirement but also a competitive advantage in human-centered industries.

In summary, the current regulatory frameworks and governance models provide an essential foundation for the ethical implementation of AI in Industry 4.0. However, sustainable progress depends on bridging the gap between principle and practice — requiring the incorporation of accountability at both technical and organizational levels, as well as the effective engagement of all actors across the value chain.

1) *Legal and Ontological Grounding for AI Governance:* In industrial ecosystems, governance should operationalize *duty of care*, accountability and traceability with meaningful human oversight, supported by risk-based assessments (e.g., DPIA under GDPR/LGPD) and auditable controls [40], [62]. We propose an ontological scaffold: {Agent, System, Decision, Evidence, Harm, Affected party, Context, Normative requirement} with relations (*produces, justifies, impacts, is-responsible-for, is-governed-by*). Coupled with documentation artifacts—*Model Cards* and *Datasheets*—and management standards (e.g., ISO/IEC 42001), this enables audit queries such as “who decided what, based on which evidence, under which constraints, and why alternative options were rejected?” [53], [54], [61].

D. Algorithmic Bias and Fairness in Socio-Technical Systems

Algorithmic systems are often perceived as objective tools, yet growing empirical evidence reveals that algorithmic bias is a pervasive and systemic issue across domains such as hiring, lending, policing, and healthcare. In industrial and operational contexts, where data-driven decision-making is increasingly central, such biases pose critical ethical, legal, and reputational risks.

Bias in AI systems typically originates from three primary sources: historical or unbalanced training data, flawed model assumptions, and systemic social inequalities that are reproduced in technical infrastructures [45]. For instance, Obermeyer et al. [46] demonstrated how a widely used healthcare algorithm in the U.S. encoded racial bias by underestimating the medical needs of Black patients based on cost-based proxies. These findings underscore the urgent need for auditing practices and fairness-aware machine learning frameworks.

Several taxonomies of fairness have emerged in the literature. Binns [21] applies insights from political philosophy

to categorize fairness notions into procedural, distributive, and relational dimensions, emphasizing the importance of contextualizing fairness depending on the system’s function and the population affected. Similarly, Barocas et al. [47] argue that technical fairness metrics (e.g., equalized odds, demographic parity) must be interpreted alongside normative and institutional considerations, especially when trade-offs are inevitable.

In industrial contexts, algorithmic bias may appear in areas such as automated recruitment, predictive maintenance, or resource allocation—where decisions affect employment, safety, and workload distribution. Holstein et al. [48] note that the lack of practical tools and organizational incentives for mitigating bias leads many teams to deprioritize fairness, especially in high-pressure environments.

To address these risks, a number of tools and practices have been proposed. One example is IBM’s AI Fairness 360 (AIF360) toolkit, which provides open-source libraries for bias detection and mitigation [49]. However, technical interventions alone are insufficient. As D’Ignazio and Klein [24] argue, “data feminism” offers a critical lens to expose how power dynamics shape data collection, labeling, and usage. They advocate for inclusive data practices and participatory design as foundational to ethical AI.

Accountability is also key to fairness. Selbst et al. [50] introduce the concept of “the portability trap”, where fairness solutions developed in one context are mistakenly assumed to be transferable across domains. This reinforces the need for localization, stakeholder engagement, and continuous monitoring.

In sum, mitigating algorithmic bias requires a multi-layered approach that encompasses technical robustness, model interpretability, inclusive data governance, and institutional accountability. For Industry 4.0 and 5.0 to evolve along ethical lines, fairness must be embedded not only in the algorithms themselves but also within the broader organizational culture and policy frameworks that guide their use.

To further illustrate these considerations, Table IV presents a summary of key ethical dilemmas associated with artificial intelligence and contrasts them with corresponding principles of responsible innovation. This comparison highlights the practical implications of these issues for the development, deployment, and governance of AI technologies across industrial settings.

Finally, to synthesize the core concepts discussed throughout this section, Fig. 2 illustrates the cyclical process of responsible innovation in artificial intelligence. The diagram emphasizes the interdependent stages — including impact assessment, stakeholder engagement, transparency, accountability, and continuous improvement — that collectively guide the ethical design and implementation of AI systems. This model reinforces the idea that ethical AI is not a static end state but an ongoing commitment throughout the system’s lifecycle.

V. CONCLUDING REMARKS ON THE LITERATURE REVIEW

The literature reviewed in this Section demonstrates a growing academic and institutional commitment to understanding and addressing the ethical implications of AI, partic-

TABLE IV
ETHICAL DILEMMAS VS. PRINCIPLES OF RESPONSIBLE INNOVATION

Ethical Dilemmas in AI	Principles of Responsible Innovation	Practical Implications
Algorithmic discrimination	Justice and Inclusion	Ensure data diversity and bias audits.
Opacity and “black-box” decision-making	Transparency and Explainability	Develop interpretable AI and communicate decisions clearly.
Diffused responsibility (who is accountable?)	Accountability	Establish clear technical and legal frameworks for liability.
Automation without human oversight	Human Autonomy and Control	Preserve human-in-the-loop systems in critical decisions.
Exploitation of personal data	Privacy and Informed Consent	Apply privacy-by-design and ensure user awareness and consent.
Profit-driven innovation vs. public good	Sustainability and Collective Interest	Align technological goals with long-term social and environmental values.

ularly in the context of Industry 4.0 and 5.0. Across all axes explored — ranging from foundational ethical principles to applied governance and algorithmic fairness — a common thread emerges: the urgent need for operationalizing ethical values within socio-technical systems.

In response to RQ1, we observed that ethical concerns surrounding AI implementation are multidimensional, encompassing transparency, bias, accountability, privacy, and human dignity. These challenges are especially pronounced in industrial environments, where automated systems interact with human workers and impact decisions with material consequences.

Addressing RQ2, recent literature reveals that although principles such as fairness, accountability, and transparency

are widely recognized, their practical application remains inconsistent. Human-centric frameworks are gaining prominence as a corrective response, calling for inclusive design processes, interpretable models, and value-sensitive system development. Notably, Industry 5.0 repositions AI from a tool of automation to a means of enhancing human well-being and creativity.

Finally, in relation to RQ3, multiple regulatory and governance proposals—ranging from the EU AI Act to localized organizational ethics frameworks — are beginning to take shape. However, the literature warns that principles without implementation mechanisms are insufficient. Sustained progress will require institutional accountability, multi-stakeholder participation, and continuous evaluation throughout the AI lifecycle.

Furthermore, the findings of this review reinforce the imperative for the Computational Intelligence community not only to pursue more efficient algorithms but also to integrate ethical considerations from the outset — an approach aligned with ethics by design. This includes the development of robust fairness metrics for machine learning, the advancement of XAI techniques applicable across a variety of computational intelligence models, and the design of mechanisms to ensure meaningful human control over autonomous systems, particularly those based on reinforcement learning or evolutionary computation.

Overall, this review highlights both the maturity and fragmentation of the field: while conceptual clarity is improving, empirical applications of ethical frameworks in real-world industrial systems remain limited. Bridging this gap is a critical task for researchers, engineers, policymakers, and organizations striving to align technological innovation with human values and social responsibility.

VI. FINAL CONSIDERATIONS

The transition toward human-centric and ethically grounded AI systems represents one of the most critical challenges and opportunities of Industry 5.0. As this article has argued, the integration of artificial intelligence into industrial processes must go beyond efficiency and performance metrics to incorporate values such as fairness, transparency, and human well-being.

Through a critical review of recent literature and ethical frameworks, we have explored how responsible innovation can be guided by normative principles, technical transparency, and inclusive governance mechanisms. However, translating these ideals into practice remains an ongoing challenge. The persistence of ethical dilemmas — such as algorithmic bias, opacity, and accountability gaps — requires not only technical innovation but also institutional and cultural transformation.

The responsible deployment of AI in Industry 4.0 contexts will depend on robust regulation, participatory design processes, and a commitment to long-term societal goals. This includes fostering interdisciplinary collaboration and engaging stakeholders across the value chain — from engineers and operators to policymakers and civil society.

Ultimately, responsible AI is not a fixed destination but a continuous process of reflection, adaptation, and co-governance. By embedding ethical principles into the design,



Fig. 2. The Responsible Innovation Cycle in Artificial Intelligence. A visual representation of key principles and feedback loops that underpin ethical and sustainable AI development

deployment, and monitoring of AI systems, we can shape a future where technology serves not just productivity, but also human dignity and sustainability.

VII. LIMITATIONS AND FUTURE RESEARCH

Although this study has thoroughly explored the ethical, regulatory, and operational challenges of implementing artificial intelligence in Industry 4.0, several important limitations must be acknowledged. Firstly, the analysis focused predominantly on normative frameworks and theoretical proposals drawn from scientific and institutional literature. The absence of specific empirical case studies limits the assessment of the practical effectiveness of these models in real industrial environments, where organizational, cultural, and technological dynamics may create tensions or deviations between proposed principles and their day-to-day application.

In addition, the adopted approach emphasized a general and cross-sectoral perspective of Industry 4.0, which may fail to capture relevant nuances of specific sectors such as the food, energy, or logistics industries, where the impacts of AI present unique characteristics. The challenges of AI adoption in developing countries were also not explored in depth, despite institutional conditions, infrastructure, and technological maturity differing significantly from those of the core nations often cited in international guidelines.

Another limitation lies in the persistent difficulty of translating widely accepted ethical principles — such as fairness, transparency, and responsibility — into concrete technical and organizational requirements. This operationalization gap, as highlighted by several authors, underscores an urgent need for methodologies that integrate ethics from system design (ethics by design) to the continuous auditing of AI systems, especially in industrial contexts where automated decisions may directly affect safety, human labor, and the environment.

In this regard, future research should prioritize three complementary fronts. The first concerns empirical investigation into how industrial organizations are implementing — or failing to implement — ethical governance mechanisms, particularly regarding algorithmic accountability and human oversight. The second involves the development of technical tools and auditable metrics that enable real-time transparency and traceability of algorithmic decisions. Lastly, interdisciplinary and participatory studies that incorporate the perspectives of workers, affected communities, engineers, and policymakers are needed to build more inclusive and adaptable co-governance models.

By addressing these gaps, it will be possible to advance not only in the construction of more trustworthy systems aligned with ethical values but also in the consolidation of an organizational culture that understands AI not merely as an enabling technology, but as a strategic factor requiring continuous ethical commitment and distributed responsibility.

Moreover, the generalizability of our findings is constrained by publication bias and the prevalence of conceptual work. Future studies should pre-register protocols, report auditable fairness and explainability metrics in production settings, and release documentation artifacts (model/data cards) and decision logs for independent scrutiny [53], [54].

VIII. ACKNOWLEDGMENTS

This work has been supported by the following Brazilian research agencies: FAPEMIG, CAPES, CNPq.

REFERENCES

- [1] F. Doshi-Velez and B. Kim, “Towards a rigorous science of interpretable machine learning,” Mar. 02, 2017, arXiv: arXiv:1702.08608. doi: 10.48550/arXiv.1702.08608.
- [2] A. Barredo Arrieta et al., “Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI,” *Information Fusion*, vol. 58, pp. 82–115, Jun. 2020, doi: 10.1016/j.inffus.2019.12.012.
- [3] E. Tjoa and C. Guan, “A Survey on Explainable Artificial Intelligence (XAI): Toward Medical XAI,” *IEEE Transactions on Neural Networks and Learning Systems*, vol. 32, no. 11, pp. 4793–4813, Nov. 2021, doi: 10.1109/TNNLS.2020.3027314.
- [4] B. D. Mittelstadt, P. Allo, M. Taddeo, S. Wachter, and L. Floridi, “The ethics of algorithms: Mapping the debate,” *Big Data & Society*, vol. 3, no. 2, p. 2053951716679679, Dec. 2016, doi: 10.1177/2053951716679679.
- [5] T. Papadakis, I. T. Christou, C. Ipektsidis, J. Soldatos, and A. Amicone, “Explainable and transparent artificial intelligence for public policymaking,” *Data & Policy*, vol. 6, p. e10, Jan. 2024, doi: 10.1017/dap.2024.3.
- [6] C. Högberg, “Stabilizing translucencies: Governing AI transparency by standardization,” *Big Data & Society*, vol. 11, no. 1, p. 20539517241234298, Mar. 2024, doi: 10.1177/20539517241234298.
- [7] Y. Zeng, E. Lu, and C. Huangfu, “Linking Artificial Intelligence principles,” Dec. 12, 2018, arXiv: arXiv:1812.04814. doi: 10.48550/arXiv.1812.04814.
- [8] H. Felzmann, E. Fosch-Villaronga, C. Lutz, and A. Tamò-Larriex, “Towards Transparency by Design for Artificial Intelligence,” *Sci Eng Ethics*, vol. 26, no. 6, pp. 3333–3361, Dec. 2020, doi: 10.1007/s11948-020-00276-4.
- [9] M. J. Page et al., “The PRISMA 2020 statement: an updated guideline for reporting systematic reviews,” Mar. 2021, doi: 10.1136/bmj.n71.
- [10] E. H. Grosse, Sgarbossa Fabio, Berlin Cecilia, and W. P. and Neumann, “Human-centric production and logistics system design and management: Transitioning from Industry 4.0 to Industry 5.0,” *International Journal of Production Research*, vol. 61, no. 22, pp. 7749–7759, Nov. 2023, doi: 10.1080/00207543.2023.2246783.
- [11] M. Tahaei, M. Constantinides, D. Quercia, and M. Muller, “A systematic literature review of human-centered, ethical, and responsible AI,” Jun. 26, 2023, arXiv: arXiv:2302.05284. doi: 10.48550/arXiv.2302.05284.
- [12] A. A. Khan et al., “Ethics of AI: A systematic literature review of principles and challenges,” in *The International Conference on Evaluation and Assessment in Software Engineering 2022*, Gothenburg Sweden: ACM, Jun. 2022, pp. 383–392. doi: 10.1145/3530019.3531329.
- [13] D. K. Gao, A. Haverly, S. Mittal, J. Wu, and J. Chen, “AI Ethics: A bibliometric analysis, critical issues, and key gaps,” *International Journal of Business Analytics*, vol. 11, no. 1, pp. 1–19, Feb. 2024, doi: 10.4018/IJBAN.338367.
- [14] L. Bibby and B. Dehe, “Defining and assessing industry 4.0 maturity levels – case of the defence sector,” *Production Planning & Control*, Sep. 2018, doi: 10.1080/09537287.2018.1503355.
- [15] R. S. Peres, X. Jia, J. Lee, K. Sun, A. W. Colombo, and J. Barata, “Industrial Artificial Intelligence in Industry 4.0 - Systematic review, challenges and outlook,” in *IEEE Access*, vol. 8, pp. 220121–220139, 2020, doi: 10.1109/ACCESS.2020.3042874.
- [16] Q. Zhu and J. Luo, “Toward artificial empathy for Human-Centered design: A framework,” May 13, 2023, arXiv: arXiv:2303.10583. doi: 10.48550/arXiv.2303.10583.
- [17] J. Morley, C. Machado, C. Burr, J. Cows, M. Taddeo, and L. Floridi, “The debate on the ethics of AI in health care: A reconstruction and critical review,” Nov. 13, 2019, Social Science Research Network, Rochester, NY: 3486518. doi: 10.2139/ssrn.3486518.
- [18] W. Samek, T. Wiegand, and K.-R. Müller, “Explainable Artificial Intelligence: Understanding, Visualizing and Interpreting Deep Learning Models,” Aug. 28, 2017, arXiv: arXiv:1708.08296. doi: 10.48550/arXiv.1708.08296.
- [19] R. Guidotti, A. Monreale, S. Ruggieri, F. Turini, F. Giannotti, and D. Pedreschi, “A Survey of Methods for Explaining Black Box Models,” *ACM Comput. Surv.*, vol. 51, no. 5, p. 93:1–93:42, Aug. 2018, doi: 10.1145/3236009.

- [20] B. Shneiderman, "Bridging the gap between ethics and practice: Guidelines for reliable, safe, and trustworthy human-centered AI systems," *ACM Trans. Interact. Intell. Syst.*, vol. 10, no. 4, pp. 1–31, Dec. 2020, doi: 10.1145/3419764.
- [21] R. Binns, "Fairness in machine learning: Lessons from political philosophy," in *Proceedings of the 1st Conference on Fairness, Accountability and Transparency*, PMLR, Jan. 2018, pp. 149–159. Accessed: May 01, 2025. [Online]. Available: <https://proceedings.mlr.press/v81/binns18a.html>.
- [22] B. Mittelstadt, "Principles alone cannot guarantee ethical AI," *Nat Mach Intell*, vol. 1, no. 11, pp. 501–507, Nov. 2019, doi: 10.1038/s42256-019-0114-4.
- [23] V. Eubanks, "Automating inequality: How high-tech tools profile, police, and punish the poor," Macmillan Publishers, Jan. 2018, ISBN: 9781250074317.
- [24] C. D'Ignazio and L. F. Klein, "Data feminism". in *Strong ideas series*. Cambridge, Massachusetts: The MIT Press, 2020, ISBN: 9780262044004.
- [25] A. Tóth, L. Nagy, R. Kennedy, B. Bohuš, J. Abonyi, and T. Ruppert, "The human-centric Industry 5.0 collaboration architecture," *MethodsX*, vol. 11, p. 102260, Dec. 2023, doi: 10.1016/j.mex.2023.102260.
- [26] D. Long, "The new renaissance: can businesses build human-centric tech?," *The Australian*, Dec. 01, 2024. Accessed: May 02, 2025. [Online]. Available: <https://www.theaustralian.com.au/business/growth-agenda/the-new-renaissance-can-businesses-build-human-centric-tech/news-story/6515c543d3ac9f622f497dd4f71e2f0b>
- [27] M. Breque, L. De Nul, A. Petridis, "Industry 5.0: Towards a sustainable, human-centric and resilient European industry," Luxembourg, LU: European Commission, Directorate-General for Research and Innovation, 2021. doi: 10.2777/308407.
- [28] X. Xu, Y. Lu, B. Vogel-Heuser, and L. Wang, "Industry 4.0 and Industry 5.0 - Inception, conception and perception," *Journal of Manufacturing Systems*, vol. 61, pp. 530–535, Oct. 2021, doi: 10.1016/j.jmsy.2021.10.006.
- [29] A. Adel, "Future of industry 5.0 in society: human-centric solutions, challenges and prospective research areas," *Journal of Cloud Computing*, vol. 11, no. 1, p. 40, Sep. 2022, doi: 10.1186/s13677-022-00314-5.
- [30] AI HLEG, "Ethics guidelines for trustworthy AI," European Commission, Apr. 2019. Accessed: May 01, 2025. [Online]. Available: <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>.
- [31] J. Morley, L. Floridi, L. Kinsey, and A. Elhalal, "From What to How: An Initial Review of Publicly Available AI Ethics Tools, Methods and Research to Translate Principles into Practices," *Sci Eng Ethics*, vol. 26, no. 4, pp. 2141–2168, Aug. 2020, doi: 10.1007/s11948-019-00165-5.
- [32] J. Fjeld, N. Achten, H. Hilligoss, A. Nagy, and M. Srikumar, "Principled Artificial Intelligence: Mapping Consensus in Ethical and Rights-Based Approaches to Principles for AI," Jan. 15, 2020, Social Science Research Network, Rochester, NY: 3518482. doi: 10.2139/ssrn.3518482.
- [33] I. Ahmed, G. Jeon, and F. Piccialli, "From Artificial Intelligence to Explainable Artificial Intelligence in Industry 4.0: A Survey on What, How, and Where," *IEEE Transactions on Industrial Informatics*, vol. 18, no. 8, pp. 5031–5042, Aug. 2022, doi: 10.1109/TII.2022.3146552.
- [34] A. K. Badhan, R. Gill, "Explainable machine learning model for Industrial 4.0," in *Machine Learning for Sustainable Manufacturing in Industry 4.0*, CRC Press, 2023.
- [35] K. Nikiforidis et al., "Enhancing transparency and trust in AI-powered manufacturing: A survey of explainable AI (XAI) applications in smart manufacturing in the era of industry 4.0/5.0," *ICT Express*, vol. 11, no. 1, pp. 135–148, Feb. 2025, doi: 10.1016/j.icte.2024.12.001.
- [36] M. T. Ribeiro, S. Singh, and C. Guestrin, "'Why Should I Trust You?': Explaining the Predictions of Any Classifier," Aug. 09, 2016, arXiv: arXiv:1602.04938. doi: 10.48550/arXiv.1602.04938.
- [37] J. M. Durán and K. R. Jongsma, "Who is afraid of black box algorithms? On the epistemological and ethical basis of trust in medical AI," *Journal of Medical Ethics*, vol. 47, no. 5, pp. 329–335, May 2021, doi: 10.1136/medethics-2020-106820.
- [38] J. Whittlestone, R. Nyrup, A. Alexandrova, and S. Cave, "The Role and Limits of Principles in AI Ethics: Towards a Focus on Tensions," in *Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society*, in AIES '19. New York, NY, USA: Association for Computing Machinery, Jan. 2019, pp. 195–200. doi: 10.1145/3306618.3314289.
- [39] European Commission, "Proposal for a Regulation laying down harmonised rules on artificial intelligence," European Commission. Accessed: May 02, 2025. [Online]. Available: <https://digital-strategy.ec.europa.eu/en/library/proposal-regulation-laying-down-harmonised-rules-artificial-intelligence>
- [40] M. Veale and F. Z. Borgesius, "Demystifying the Draft EU Artificial Intelligence Act," *Computer Law Review International*, vol. 22, no. 4, pp. 97–112, Aug. 2021, doi: 10.9785/crl-2021-220402.
- [41] A. F. T. Winfield and M. Jirotko, "Ethical governance is essential to building trust in robotics and artificial intelligence systems," *Philos Trans A Math Phys Eng Sci*, vol. 376, no. 2133, p. 20180085, Oct. 2018, doi: 10.1098/rsta.2018.0085.
- [42] M. Ashok, R. Madan, A. Joha, and U. Sivarajah, "Ethical framework for Artificial Intelligence and Digital technologies," *International Journal of Information Management*, vol. 62, p. 102433, Feb. 2022, doi: 10.1016/j.ijinfomgt.2021.102433.
- [43] V. Dignum, "Taking Responsibility," in *Responsible Artificial Intelligence: How to Develop and Use AI in a Responsible Way*, V. Dignum, Ed., Cham: Springer International Publishing, 2019, pp. 47–69. doi: 10.1007/978-3-030-30371-6_4.
- [44] A. Z. Coelho et al., "Governança da inteligência artificial em organizações: framework para comitês de ética em IA: versão 1.0," CEPI FGV Direito SP, May 2023. Accessed: May 02, 2025. [Online]. Available: <https://hdl.handle.net/10438/33736>
- [45] N. Mehrabi, F. Morstatter, N. Saxena, K. Lerman, and A. Galstyan, "A Survey on Bias and Fairness in Machine Learning," *ACM Comput. Surv.*, vol. 54, no. 6, p. 115:1–115:35, Jul. 2021, doi: 10.1145/3457607.
- [46] Z. Obermeyer, B. Powers, C. Vogeli, and S. Mullainathan, "Dissecting racial bias in an algorithm used to manage the health of populations," *Science*, vol. 366, no. 6464, pp. 447–453, Oct. 2019, doi: 10.1126/science.aax2342.
- [47] S. Barocas, M. Hardt, and A. Narayanan, "Fairness and machine learning," 2023. Accessed: Jan 08, 2025. [Online]. Available: <https://fairmlbook.org/>
- [48] K. Holstein, J. Wortman Vaughan, H. Daumé, M. Dudik, and H. Wallach, "Improving Fairness in Machine Learning Systems: What Do Industry Practitioners Need?," in *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems*, in CHI '19. New York, NY, USA: Association for Computing Machinery, Maio 2019, pp. 1–16. doi: 10.1145/3290605.3300830.
- [49] R. K. E. Bellamy et al., "AI Fairness 360: An Extensible Toolkit for Detecting, Understanding, and Mitigating Unwanted Algorithmic Bias," Oct. 03, 2018, arXiv: arXiv:1810.01943. doi: 10.48550/arXiv.1810.01943.
- [50] A. D. Selbst, D. Boyd, S. A. Friedler, S. Venkatasubramanian, and J. Vertesi, "Fairness and Abstraction in Sociotechnical Systems," in *Proceedings of the Conference on Fairness, Accountability, and Transparency*, in FAT* '19. New York, NY, USA: Association for Computing Machinery, Jan. 2019, pp. 59–68. doi: 10.1145/3287560.3287598.
- [51] Q. N. Hong et al., "Mixed Methods Appraisal Tool (MMAT) version 2018: user guide," 2018. [Online]. Available: <http://mixedmethodsappraisaltoolpublic.pbworks.com>
- [52] Q. N. Hong, V. Fàbregues, and P. Bartlett, "The Mixed Methods Appraisal Tool (MMAT) version 2022," 2022. [Online]. Available: <http://mixedmethodsappraisaltoolpublic.pbworks.com>
- [53] M. Mitchell et al., "Model Cards for Model Reporting," in *FAT**, 2019, pp. 220–229. doi: 10.1145/3287560.3287596.
- [54] T. Gebru et al., "Datasheets for Datasets," *Communications of the ACM*, 64(12), pp. 86–92, 2021. doi: 10.1145/3458723.
- [55] M. Hardt, E. Price, and N. Srebro, "Equality of Opportunity in Supervised Learning," in *NeurIPS*, 2016, pp. 3315–3323.
- [56] S. Wachter, B. Mittelstadt, and C. Russell, "Counterfactual Explanations without Opening the Black Box: Automated Decisions and the GDPR," *Harvard Journal of Law & Technology*, 31(2), 2018.
- [57] M. J. Kusner, J. Loftus, C. Russell, and R. Silva, "Counterfactual Fairness," in *NeurIPS*, 2017, pp. 4066–4076.
- [58] S. M. Lundberg and S.-I. Lee, "A Unified Approach to Interpreting Model Predictions," in *NeurIPS*, 2017, pp. 4765–4774.
- [59] A. N. Angelopoulos and S. Bates, "A Gentle Introduction to Conformal Prediction and Distribution-Free Uncertainty Quantification," *arXiv:2107.07511*, 2021.
- [60] M. Abadi et al., "Deep Learning with Differential Privacy," in *ACM CCS*, 2016, pp. 308–318. doi: 10.1145/2976749.2978318.
- [61] ISO/IEC 42001:2023, "Artificial intelligence — Management system," International Organization for Standardization, 2023.
- [62] Brasil, "Lei nº 13.709, de 14 de agosto de 2018 (Lei Geral de Proteção de Dados Pessoais)," 2018. [Online]. Available: <https://www.planalto.gov.br/>

ARTIGO II: A URGÊNCIA ÉTICA DA IA CENTRADA NO HUMANO NA ERA DA AUTOMAÇÃO DO TRABALHO

Trabalho submetido e aprovado no “16th IEEE/IAS International Conference on Industry Applications - INDUSCON 2025”, em maio de 2025.

A Urgência Ética da IA Centrada no Humano na Era da Automação do Trabalho

1st Julio C. S. Lima
Programa de Pós-Graduação em
Engenharia de Sistemas e Automação
Universidade Federal de Lavras
Lavras/MG, Brasil
lima.julio.cs@gmail.com

2nd Felipe O. Silva
Departamento de Automática
Universidade Federal de Lavras
Lavras/MG, Brasil
felipe.oliveira@ufla.br

3rd Danton Diego Ferreira
Departamento de Automática
Universidade Federal de Lavras
Lavras/MG, Brasil
danton@ufla.br

Abstract—A Inteligência Artificial (IA) está transformando rapidamente os contextos industriais e sociais, proporcionando ganhos inéditos de eficiência. No entanto, essa transformação traz preocupações éticas urgentes, especialmente no que se refere à substituição de trabalhadores e à obsolescência de competências, agravadas pela ausência de mecanismos adequados de apoio e requalificação. Este artigo tem como objetivo analisar criticamente as implicações éticas da automação do trabalho baseada em IA, com ênfase na adoção de uma IA Centrada no Humano (HCAI) como paradigma mais responsável e sustentável. As principais contribuições incluem: (i) uma revisão sistemática dos conceitos de Indústria 4.0, Indústria 5.0 e HCAI sob uma perspectiva ética; (ii) a articulação dos desafios da automação à luz dos princípios de justiça social; e (iii) a proposição de diretrizes para políticas públicas, práticas educacionais e responsabilidade corporativa voltadas à promoção de uma transição tecnológica justa e inclusiva. Como resultado, o estudo identifica lacunas críticas nas políticas atuais de automação e destaca a urgência ética de uma mudança de modelos orientados apenas pela eficiência para inovações centradas no ser humano. O artigo conclui com um conjunto de recomendações normativas e práticas que reforçam a importância da equidade, sustentabilidade e resiliência na configuração do futuro do trabalho.

Index Terms—Inteligência Artificial, Ética da Inteligência Artificial, IA Centrada no Humano, Indústria 5.0, Desemprego Tecnológico, Requalificação Profissional.

I. INTRODUÇÃO: O PARADOXO DA IA - PROGRESSO E PERIGO HUMANO

A Inteligência Artificial (IA) permeia cada vez mais a sociedade contemporânea, otimizando processos, impulsionando a inovação e redefinindo indústrias inteiras [1]. Desde sistemas de recomendação a diagnósticos médicos assistidos e automação industrial avançada, a capacidade da IA de processar dados e realizar tarefas complexas oferece um potencial transformador inegável [2]. Contudo, por trás das promessas de eficiência e progresso, emerge uma tensão profunda com implicações humanas significativas: o impacto da automação impulsionada pela IA no mercado de trabalho [3].

Estudos influentes alertam para a alta suscetibilidade de um grande número de ocupações à informatização [4]–[6]. A capacidade da IA de automatizar não apenas tarefas manuais repetitivas, mas também tarefas cognitivas complexas, ameaça deslocar trabalhadores em uma escala potencialmente sem precedentes [7], [8], levantando o espectro do desemprego tecnológico e da crescente desigualdade [9].

Este cenário transcende a mera análise econômica, configurando-se como um dos desafios éticos mais prementes de nossa era. A perspectiva de milhões de trabalhadores tornarem-se obsoletos, enfrentando insegurança e perda de propósito, exige uma reflexão sobre a justiça, a equidade e a dignidade humana no desenvolvimento tecnológico [10]. Agravando essa preocupação está a constatação de que as estruturas atuais de apoio – sejam políticas públicas, sistemas educacionais ou iniciativas corporativas – mostram-se frequentemente insuficientes ou inadequadas para facilitar a massiva transição profissional que se avizinha [11]. Essa lacuna não é apenas uma questão prática, mas uma falha ética que clama por abordagens alternativas.

Diante desse cenário de transformação impulsionado pela IA e seus impactos no mercado de trabalho, torna-se essencial compreender os conceitos que fundamentam essa nova realidade tecnológica e social. Para aprofundar a análise, a Seção III propõe a contextualização da transformação em curso, abordando os princípios centrais da Indústria 4.0 e da emergente Indústria 5.0, suas dinâmicas de interação homem-máquina — com ênfase nos paradigmas “Humano-no-Ciclo” (*Human-In-The-Loop* - HITL) e “IA Centrada no Humano” (*Human-Centric AI* - HCAI) —, e o papel da ética como eixo transversal nesse processo. Essa abordagem permite não apenas situar a discussão em seu momento histórico e tecnológico, mas também evidenciar a importância crítica do presente estudo na compreensão das oportunidades e desafios que marcam esta fase de convergência entre inovação, trabalho e responsabilidade social.

Este artigo tem como objetivo principal analisar criticamente os impactos éticos da automação do trabalho impulsionada pela IA, com ênfase na proposta de uma abordagem centrada no ser humano. As contribuições deste trabalho à comunidade acadêmica incluem: (i) a sistematização dos conceitos de Indústria 4.0, Indústria 5.0 e HCAI sob uma perspectiva ética; (ii) a articulação entre os desafios distributivos da automação e os princípios de justiça social; e (iii) a proposição de diretrizes para políticas públicas, práticas educacionais e responsabilidade corporativa voltadas à promoção de uma transição tecnológica justa e inclusiva.

II. METODOLOGIA

A presente pesquisa adota uma abordagem qualitativa e exploratória, com foco em análise crítica conceitual e

normativo-ética. Utiliza-se uma combinação de revisão narrativa com sistematização temática, ancorada na literatura científica sobre IA, automação do trabalho, Indústria 5.0, ética tecnológica e HCAI.

A construção do argumento seguiu os seguintes passos metodológicos:

- **Definição do problema e delimitação do escopo:** mapeamento do paradoxo entre automação e centralidade humana na Indústria 4.0 e 5.0;
- **Revisão crítica da literatura:** levantamento de publicações científicas recentes (2015–2025), com ênfase em documentos da OECD, IEEE, Springer Nature Link, ACM Digital Library, MDPI, arXiv e periódicos revisados por pares;
- **Análise ética normativa:** identificação de implicações éticas centrais relacionadas ao deslocamento de trabalho, desigualdade e justiça social;
- **Proposição de diretrizes:** construção argumentativa de soluções centradas no paradigma HCAI, com desdobramentos em políticas públicas, educação e responsabilidade corporativa.

Trata-se, portanto, de uma pesquisa teórico-analítica cujo objetivo é fornecer um modelo interpretativo e diretivo para guiar a inovação tecnológica de forma responsável.

Problema de Pesquisa: Como a adoção acrítica de IA na automação do trabalho pode comprometer a justiça social, e quais princípios normativos devem nortear uma transição tecnológica mais ética e centrada no humano?

III. CONTEXTUALIZANDO A TRANSFORMAÇÃO: CONCEITOS FUNDAMENTAIS

A rápida evolução da IA insere-se em um contexto industrial e tecnológico caracterizado por transformações profundas, como as propostas pelas Indústrias 4.0 e 5.0. Enquanto a Indústria 4.0 enfatiza a automação e a digitalização dos processos produtivos [12], a Indústria 5.0 propõe uma abordagem mais humanizada, centrada na colaboração entre humanos e máquinas [13]. Nesse cenário, abordagens como HITL e HCAI ganham destaque, promovendo interações mais éticas e eficazes entre operadores e sistemas inteligentes [14], [15]. A literatura recente destaca também a necessidade de reformular práticas de *design* de sistemas de IA para assegurar sua orientação humanocêntrica [16]. Compreender esses conceitos e suas inter-relações é fundamental para analisar as implicações éticas da automação e fundamentar a busca por uma inovação responsável e sustentável.

A. Indústria 4.0 vs. Indústria 5.0: Da Eficiência Tecnológica ao Valor Humano-Social

A Quarta Revolução Industrial, consolidada sob o termo Indústria 4.0, marcou a última década com uma profunda transformação nos sistemas produtivos, originada na Alemanha [17]. Caracterizada pela digitalização massiva e pela interconexão viabilizada por tecnologias como Sistemas Ciber-Físicos (CPS), Internet das Coisas (IoT) e *Big Data*, seu foco primordial reside na automação, otimização de processos, e no aumento da eficiência e flexibilidade da produção [12], [13], [17], [18]. Esta fase é amplamente reconhecida como uma transformação impulsionada primariamente pela tecnologia [17]. Contudo, apesar dos inegáveis

avanços em produtividade, a forte ênfase da Indústria 4.0 na automação tecnológica suscitou crescentes preocupações éticas e sociais, notadamente sobre a potencial substituição do trabalho humano e suas consequências.

Diante dessas preocupações e da necessidade de alinhar o desenvolvimento industrial com valores sociais mais amplos, emergiu o conceito de Indústria 5.0, formalmente apresentado pela Comissão Europeia [17], [19]. Longe de ser uma mera sucessora cronológica ou uma substituta, a Indústria 5.0 é proposta como um complemento e uma extensão da Indústria 4.0 [17], [20], reorientando o foco para uma abordagem impulsionada por valores. Seus três pilares centrais – humanocentrismo, que reposiciona o bem-estar e as necessidades do trabalhador no coração do processo produtivo; sustentabilidade, que exige respeito aos limites planetários e promove a economia circular; e resiliência, que busca robustez e adaptabilidade frente a crises e disrupções [13], [17], [19] – sinalizam um reconhecimento fundamental. A transição para a Indústria 5.0 reflete, portanto, a crescente convicção de que a inovação tecnológica não pode se dissociar de objetivos sociais e ambientais, devendo, ao contrário, promover uma colaboração mais sinérgica e harmoniosa entre humanos e máquinas [18], [21]–[23], um ponto central para a discussão ética sobre o futuro do trabalho neste artigo.

B. Interação Humano-Máquina: HITL vs. HCAI

A interação entre humanos e sistemas de IA é um elemento crucial no desenvolvimento e operação dessas tecnologias, manifestando-se em diferentes níveis e com distintos propósitos. No debate sobre o futuro do trabalho e a automação, é fundamental distinguir entre conceitos como HITL e HCAI, pois eles representam filosofias e práticas distintas de interação humano-máquina.

HITL refere-se tipicamente a arquiteturas de sistema onde um humano intervém ativamente em etapas específicas do ciclo operacional ou de aprendizado da IA [24]. Essa intervenção é muitas vezes necessária para tarefas como rotulagem de dados complexos, validação de previsões de baixa confiança do modelo, correção de erros algorítmicos, ou gerenciamento de exceções e casos de borda que a IA não consegue tratar autonomamente [24]. O objetivo primário do HITL é, frequentemente, instrumental: melhorar a precisão, a robustez ou a segurança do próprio sistema de IA, utilizando o julgamento ou a capacidade humana como um componente funcional para superar as limitações da máquina [25]–[27]. A perspectiva pode ser vista como sistema-cêntrica, onde o humano serve ao aprimoramento da tecnologia.

Em contraste, a HCAI representa uma abordagem de *design* e uma filosofia de desenvolvimento mais abrangente [28]. A HCAI não vê o humano apenas como um componente para otimizar a máquina, mas posiciona as necessidades, valores, capacidades e o bem-estar humano como o ponto central e o objetivo final do desenvolvimento da IA [28]. O propósito transcende a mera funcionalidade; busca-se criar sistemas de IA que aumentem as capacidades humanas, preservem a autonomia e a dignidade do usuário, garantam transparência e segurança, e contribuam para objetivos sociais positivos e equitativos [29]–[31]. A HCAI advoga por uma parceria sinérgica onde a tecnologia amplifica

o potencial humano e respeita seus limites e valores [17], [32].

A distinção é crucial: enquanto a HCAI frequentemente requer mecanismos de HITL bem projetados para garantir o controle humano significativo e a agência do usuário [29], [33], a mera presença de um humano no ciclo não garante que um sistema seja verdadeiramente humanocêntrico. Um sistema HITL pode tratar o humano como uma mera ferramenta ou engrenagem, otimizando a eficiência sem considerar seu bem-estar ou autonomia [24]. No contexto deste artigo, que aborda os desafios éticos da automação no trabalho, a HCAI emerge como a abordagem normativa desejada para guiar o desenvolvimento de IA, visando mitigar riscos de deslocamento e promover um futuro de trabalho mais colaborativo e digno, onde o HITL é implementado como um meio para alcançar esses fins humanísticos [17].

Para cristalizar as diferenças fundamentais entre essas duas abordagens de interação humano-máquina, a Tabela I resume suas características distintivas quanto ao foco principal, objetivo primário, nível de aplicação predominante, ênfase conceitual e um exemplo prático ilustrativo.

TABELA I
COMPARATIVO ENTRE HITL E HCAI.

Aspecto	HITL	HCAI
Foco	Participação humana funcional no ciclo da IA	Valores humanos, bem-estar, ética, agência humana
Objetivo Primário	Melhorar <i>performance</i> / robustez / segurança da IA	Amplificar capacidades humanas; IA justa, segura, responsável
Nível de Aplicação	Implementação operacional/técnica	Filosofia de <i>design</i> ; Propósito e impacto do sistema
Ênfase	Supervisão/intervenção técnica ou para tarefas específicas	Princípios éticos, sociais, controle humano significativo
Exemplo Prático	Moderador humano revisando conteúdo sinalizado por IA	IA que adapta interface e tarefas às necessidades do trabalhador

C. Ética: A Bússola para a Inovação Responsável

A ética, em seu sentido fundamental, refere-se ao conjunto de princípios morais que orientam o comportamento humano e a condução de atividades. No contexto tecnológico, especialmente no desenvolvimento e aplicação da IA, a ética torna-se essencial para avaliar ações e artefatos em termos de certo e errado, justo e injusto, considerando seu impacto sobre indivíduos, sociedade e meio ambiente [34].

A ética na IA não é um campo acessório, mas uma necessidade intrínseca para garantir que os sistemas desenvolvidos respeitem valores fundamentais como justiça, transparência e responsabilidade. Estudos recentes destacam preocupações sobre a perpetuação de vieses e discriminações por sistemas de IA, especialmente em processos de recrutamento e seleção, onde algoritmos podem reforçar desigualdades existentes se não forem cuidadosamente projetados e monitorados [35].

Além disso, a proteção da privacidade e da autonomia dos usuários é uma consideração ética central. A coleta e o uso de dados pessoais por sistemas de IA levantam questões sobre

consentimento e controle, exigindo abordagens que garantam a transparência dos processos de decisão e a segurança das informações [36], [37].

Crucialmente, a consideração das consequências sociais e econômicas da IA, como o impacto no emprego, é fundamental. A automação impulsionada pela IA pode levar à substituição de empregos humanos, exigindo uma análise ética sobre como mitigar efeitos negativos e promover uma transição justa para os trabalhadores afetados [34], [38].

Portanto, integrar considerações éticas no desenvolvimento e aplicação da IA é vital para assegurar que essa tecnologia contribua positivamente para a sociedade, promovendo equidade, respeito aos direitos humanos e bem-estar coletivo.

D. A Interconexão no Momento Tecnológico Atual e a Relevância para o Artigo

O panorama tecnológico contemporâneo é definido por uma convergência crítica: de um lado, o potencial disruptivo da automação, exponenciado pela Indústria 4.0 e pela IA; de outro, a intensificação das preocupações éticas sobre as profundas consequências humanas e sociais dessa transformação. Nesse contexto de tensão inerente, emergem a visão da Indústria 5.0 [17] e a filosofia da HCAI [29] não apenas como respostas, mas como propostas para um novo paradigma de inovação. Tais abordagens buscam ativamente reconciliar o avanço tecnológico com a centralidade dos valores humanos, a dignidade e o bem-estar [18], [29], marcando uma evolução em relação aos modelos puramente tecnocêntricos anteriores [17].

Compreender essa dinâmica intrincada é fundamental para a análise central deste artigo. A automação, impulsionada pela Indústria 4.0, apresenta dinâmicas complexas de deslocamento e potencial recriação de tarefas [39], levantando desafios éticos prementes sobre o futuro do trabalho e a dignidade humana [40]. A Indústria 4.0, portanto, representa o vetor tecnológico que catalisa a automação e, conseqüentemente, gera a preocupação ética sobre o futuro do trabalho. A Ética, por sua vez, fornece o arcabouço normativo e a lente crítica indispensáveis para avaliar as ramificações dessa automação (cf. [40]) – a negligência ou a falta de suporte adequado aos trabalhadores afetados configura-se, inequivocamente, como uma falha ética.

Em contrapartida, a Indústria 5.0 e a HCAI oferecem um horizonte alternativo. A Indústria 5.0 é explicitamente proposta como uma “solução centrada no humano” [18], focando na colaboração humano-máquina, na resiliência e na sustentabilidade social [17], [18]. Similarmente, a HCAI advoga por sistemas que amplifiquem as capacidades humanas, garantindo controle e compreensão por parte dos usuários [29]. Estratégias como os sistemas HITL podem servir como mecanismos práticos para concretizar essa visão humanocêntrica, assegurando a agência, o controle e o envolvimento humano nos processos automatizados, alinhando-se aos princípios da HCAI [29].

Portanto, a inter-relação entre esses conceitos – a força disruptiva (Indústria 4.0/IA, analisada economicamente por Acemoglu e Restrepo [39]), a avaliação crítica (Ética, como discutido por Breque et al. [40]) e a proposta de solução humanocêntrica (Indústria 5.0 e HCAI, conforme [17], [18], [29]) – delineia o complexo cenário no qual a automação

do trabalho deve ser navegada. Essa interação justifica e sublinha a urgência ética, tema central deste artigo, de se adotar e implementar abordagens que priorizem, de forma equilibrada e integrada, tanto a inovação tecnológica quanto a dignidade, o desenvolvimento e a prosperidade da força de trabalho humana.

IV. A DIMENSÃO ÉTICA DA AUTOMAÇÃO DO TRABALHO: DESLOCAMENTO, OBSOLESCÊNCIA E A RESPONSABILIDADE PELA TRANSIÇÃO

A introdução da IA nos processos produtivos não é um mero avanço técnico; ela reconfigura profundamente o cenário do trabalho, trazendo consigo implicações éticas inescapáveis que demandam análise crítica e ação responsável [31], [41]. As consequências distributivas e sociais da automação impõem um questionamento sobre justiça, equidade e dignidade humana, exigindo uma abordagem de inovação responsável. A análise ética do impacto da IA no emprego desvela uma complexa teia de desafios interconectados [2], [42], que se manifestam em diversas dimensões críticas.

Uma das preocupações éticas mais imediatas advém do potencial de deslocamento de trabalhadores em larga escala. Embora a automação possa criar novas funções, a história e a análise econômica sugerem que essa transição não é automática nem equitativa [2], [42]. Não há garantia de que a criação de novos postos compense as perdas em todos os setores ou que sejam acessíveis àqueles cujas ocupações foram suplantadas, podendo resultar em desemprego estrutural. Crucialmente, esse processo tende a exacerbar as desigualdades socioeconômicas preexistentes. Estudos indicam que a automação amplia desigualdades ao concentrar os ganhos entre elites qualificadas, enquanto custos sociais recaem sobre grupos vulneráveis, agravados por vieses algorítmicos no ambiente de trabalho [9], [43], [44], configurando um cenário de potencial injustiça distributiva [9], [31], [43]. Além disso, é importante destacar que os efeitos da automação impulsionada pela IA não são uniformes: as consequências podem ser amplificadas entre grupos já historicamente marginalizados, como mulheres, pessoas negras, trabalhadores mais velhos ou indivíduos de regiões periféricas e países em desenvolvimento. Estudos indicam que algoritmos aplicados a processos de seleção, produtividade e recomendação frequentemente reproduzem — e até amplificam — desigualdades estruturais presentes na sociedade [35], [44], [45]. Ignorar essas camadas interseccionais pode comprometer a justiça da transição tecnológica, tornando ainda mais urgente a adoção de abordagens éticas sensíveis à diversidade.

Complementarmente, a velocidade da inovação em IA acelera a obsolescência de conjuntos de habilidades que antes garantiam estabilidade profissional [2], [42]. Isso impõe a necessidade premente de aprendizado contínuo e adaptação (*lifelong learning*), um processo dificultado pela distribuição desigual de recursos, tempo e acesso a oportunidades de requalificação eficazes [9], [46]. A perda do emprego afeta não só a renda, mas também a identidade, a autoestima e o senso de propósito social [47], [48].

Nesse contexto, talvez o ponto mais crítico da perspectiva ética seja a inadequação da resposta coletiva — por parte

de governos, setores empresariais e instituições educacionais — a essa transformação disruptiva. A ausência de políticas públicas robustas para requalificação, segurança social e partilha de ganhos da automação configura uma grave lacuna de suporte [9]. Políticas públicas específicas são necessárias para mitigar os efeitos da IA sobre a polarização do emprego e a desigualdade [9]. A relutância de parte do setor corporativo em assumir responsabilidade significativa pela transição de seus trabalhadores também levanta questões éticas sobre governança corporativa e inovação responsável [3], [31], [34], [41]. Argumenta-se que negligenciar a gestão ativa dessa transição, cujos impactos podem variar significativamente entre países desenvolvidos e em desenvolvimento [49], representa uma abdicação da responsabilidade ética fundamental de zelar pelo bem-estar coletivo e pela justiça social [49]. A inação ou o suporte insuficiente não são meras falhas operacionais, mas refletem uma questão de prioridades e valores éticos.

Em suma, a automação via IA coloca desafios éticos que transcendem a eficiência tecnológica [31]. Enfrentá-los exige uma governança proativa e humanocêntrica, fundamentada em princípios de justiça e responsabilidade [3], [41], que assegure uma transição equitativa [9] e minimize os custos humanos da transformação do trabalho [42], [47].

V. IA CENTRADA NO HUMANO: UM PARADIGMA PARA O TRABALHO COLABORATIVO

Diante dos desafios éticos impostos pela automação irrestrita, emerge a necessidade de um paradigma alternativo: a HCAI. Esta abordagem não nega o potencial da IA, mas busca direcioná-la de forma a complementar e potencializar as capacidades humanas, em vez de meramente substituí-las, promovendo um futuro de trabalho mais colaborativo e digno [50], [51].

A HCAI fundamenta-se em princípios que colocam a agência, o bem-estar e o desenvolvimento humano no centro do *design* e implementação de sistemas de IA [51]. Um princípio chave é o aumento (*augmentation*), que visa a criação de ferramentas de IA que amplifiquem a inteligência, a criatividade e a produtividade humanas. Em vez de automatizar completamente uma função, a IA de aumento permite que as pessoas realizem tarefas complexas com maior eficiência, como demonstrado em aplicações que vão desde assistência ao diagnóstico médico até o auxílio em processos criativos complexos [52]. Outro pilar essencial é a colaboração (*collaboration*), que foca em projetar sistemas onde humanos e IA trabalham sinergicamente, aproveitando as forças distintas de cada um [32]. Essa colaboração pode ocorrer em diversos níveis, desde a análise conjunta de dados até a interação física com robôs colaborativos (*cobots*) no chão de fábrica, um conceito central para a Indústria 5.0 [17], [53]. No entanto, é crucial abordar essa colaboração de forma crítica, considerando também potenciais consequências não intencionais para a segurança e autonomia do trabalhador [54].

Para que essa parceria seja ética e eficaz, garantir o controle humano significativo (*meaningful human control*) é indispensável. Isso implica assegurar que os humanos mantenham a capacidade de supervisionar, intervir e, em última instância, direcionar as ações e decisões dos sistemas de IA,

especialmente em contextos críticos, prevenindo a diluição da responsabilidade [50], [55]. Complementarmente, a HCAI exige uma atenção primordial à segurança e bem-estar do trabalhador. Isso envolve não apenas a segurança física na interação com máquinas, mas também a consideração de fatores ergonômicos e cognitivos para evitar sobrecarga e promover um ambiente de trabalho saudável e motivador, como preconizado pelas revisões sobre bem-estar e fatores humanos na Indústria 5.0 [53], [56].

Na prática, esses princípios se materializam em diversas implementações que já começam a moldar o ambiente de trabalho. As ferramentas de aumento incluem desde *softwares* que auxiliam radiologistas na interpretação de imagens até plataformas de IA generativa que colaboram com profissionais criativos [52]. Os *cobots* na indústria são exemplos tangíveis da colaboração física, assumindo tarefas repetitivas ou extenuantes ao lado de operadores humanos que se concentram em aspectos que exigem maior destreza ou julgamento [17]. Além disso, um campo promissor é o uso da própria IA para capacitação, com o desenvolvimento de plataformas de *e-learning* adaptativas que podem identificar lacunas de competência individuais e sugerir percursos personalizados de requalificação, facilitando a adaptação contínua exigida pela evolução tecnológica.

Adotar uma abordagem HCAI, portanto, não significa frear a inovação. Pelo contrário, representa um esforço consciente para direcionar o imenso potencial tecnológico da IA para um caminho que reconheça e valorize o papel insubstituível do trabalho humano. Trata-se de buscar ativamente não apenas mitigar os custos sociais da automação, mas também de cocriar um futuro do trabalho onde a tecnologia sirva como uma força para o empoderamento, a segurança e o bem-estar, maximizando os benefícios compartilhados para toda a sociedade [50], [51].

VI. O ECOSISTEMA NECESSÁRIO: POLÍTICAS, EDUCAÇÃO E RESPONSABILIDADE CORPORATIVA

A transição para um futuro de trabalho onde a IA atue de forma mais colaborativa e menos substitutiva, conforme advoga a abordagem da HCAI, não ocorrerá espontaneamente apenas pela boa vontade dos desenvolvedores ou pelas forças de mercado. A implementação bem-sucedida da HCAI e, fundamentalmente, a mitigação justa dos impactos negativos da automação sobre a força de trabalho exigem um esforço concertado e a construção de um ecossistema de apoio robusto. Este ecossistema deve preencher a lacuna de suporte atualmente existente, envolvendo ações coordenadas em múltiplas frentes: políticas públicas, sistemas educacionais e a própria responsabilidade ética das corporações.

O papel do Estado é crucial na orquestração dessa complexa transição [47], necessitando de políticas públicas ativas e adaptativas que transcendam medidas paliativas. É fundamental a criação de mecanismos de incentivo, como benefícios fiscais ou subsídios diretos, para empresas que comprovadamente adotam tecnologias de aumento (*augmentation*) em vez de substituição pura, e que investem significativamente na requalificação e realocação de seus funcionários impactados pela automação. Paralelamente, a modernização da rede de segurança social é imperativa; modelos tradicionais de seguro-desemprego podem ser insuficientes diante

de deslocamentos em larga escala, tornando necessário explorar alternativas como a Renda Básica Universal (RBU) ou a criação de contas individuais de aprendizagem [57], garantindo um suporte mínimo durante períodos de transição ou requalificação. Além disso, a governança da IA deve incluir regulações que exijam avaliações prévias de impacto no emprego para grandes projetos de automação, possivelmente atreladas a planos de mitigação e transição justa para os trabalhadores afetados [31], [58].

Complementarmente, as reformas educacionais estruturais são indispensáveis. O sistema educacional, em todos os níveis, precisa se adaptar rapidamente para preparar os cidadãos para um futuro onde a colaboração com a IA será a norma [59]. Isso implica ir além do ensino de habilidades técnicas digitais, enfatizando fortemente competências intrinsecamente humanas e complementares à IA, como pensamento crítico, resolução de problemas complexos, criatividade, inteligência socioemocional e adaptabilidade [42], [60]. Uma cultura de aprendizado ao longo da vida (*lifelong learning*) deve ser fomentada e facilitada [59], com acesso contínuo a oportunidades de atualização e aquisição de novas competências relevantes para um mercado de trabalho em constante fluxo.

Por fim, a responsabilidade social corporativa (RSC) na era da IA assume uma dimensão crítica. As empresas que desenvolvem e implementam IA detêm uma responsabilidade ética que ultrapassa a mera maximização do lucro [58]. Isso se traduz em ações concretas como priorizar o investimento no capital humano, oferecendo requalificação e desenvolvimento interno como primeira opção antes de recorrer a demissões em massa [57]. A transparência e o diálogo com os funcionários sobre planos de automação e seus impactos são essenciais para construir confiança e permitir um planejamento conjunto. Adotar um i ético e participativo, envolvendo os próprios trabalhadores na concepção de sistemas de IA que afetarão seus postos de trabalho, pode levar a soluções mais eficazes e humanocêntricas. Finalmente, a formação de parcerias estratégicas com governos e instituições educacionais é vital para cocriar programas de treinamento relevantes e eficazes, alinhados às necessidades futuras do mercado e às trajetórias de carreira dos trabalhadores [60].

Somente através da construção ativa dessas pontes – políticas públicas visionárias, um sistema educacional adaptado e uma conduta corporativa eticamente responsável [31] – será possível garantir que os benefícios da IA sejam amplamente compartilhados e que a transição tecnológica inerente à automação não aprofunde as fraturas sociais, mas sim contribua para um futuro de trabalho mais justo e próspero para todos.

VII. DISCUSSÃO DOS RESULTADOS E IMPLICAÇÕES ÉTICAS

A transição para um futuro do trabalho moldado pela IA apresenta um cenário complexo, repleto de promessas e desafios intrincados. Embora a adoção de uma abordagem de HCAI, como explorado na Seção anterior, ofereça um paradigma ético e produtivo promissor apoiado por políticas robustas, sua implementação não é isenta de obstáculos significativos. É crucial reconhecer que a priorização da colaboração humano-máquina e do bem-estar humano, em

detrimento da automação total focada apenas na eficiência, pode implicar em custos iniciais mais elevados e maior complexidade de integração no curto prazo. Como argumentam [28], projetar sistemas HCAI eficazes exige um investimento considerável em *design* de interação, treinamento e adaptação organizacional, que pode parecer menos atraente financeiramente em comparação com soluções de automação “chave na mão”. Ademais, mesmo sob o paradigma HCAI, a reconfiguração de processos produtivos inevitavelmente levará à obsolescência de certas funções e tarefas, sublinhando a necessidade contínua e urgente de mecanismos de apoio à transição justa para os trabalhadores afetados, como programas de requalificação e redes de segurança social robustas [61].

Contudo, a narrativa da HCAI transcende a mera mitigação de riscos, abrindo um vasto horizonte para a criação de novos papéis e profissões intrinsecamente ligados à interação e gestão da tecnologia. A literatura recente destaca o surgimento de funções como treinadores de IA, curadores de dados, especialistas em ética algorítmica, *designers* de experiência de interação humano-máquina e gestores de sistemas colaborativos humano-robô [62]. Essas novas oportunidades, focadas em habilidades humanas como criatividade, pensamento crítico, inteligência emocional e supervisão ética, representam um potencial significativo para enriquecer o trabalho humano. A questão fundamental que emerge, no entanto, e que ressoa com a dimensão ética discutida anteriormente, reside em como garantir que os benefícios dessas novas avenidas profissionais sejam distribuídos equitativamente, tornando-se acessíveis especialmente àqueles cujos empregos tradicionais foram deslocados pela automação [63]. A ausência de estratégias inclusivas pode exacerbar as desigualdades existentes, criando uma clivagem ainda maior na força de trabalho.

Essa tensão entre os desafios da transição e as oportunidades emergentes é frequentemente exacerbada pela resistência à mudança. Por um lado, empresas focadas em métricas de eficiência de curto prazo e maximização de lucros podem relutar em adotar modelos HCAI que demandam investimentos iniciais maiores ou uma reavaliação de processos estabelecidos [64]. Por outro lado, a força de trabalho, confrontada com a incerteza e a ansiedade sobre a segurança do emprego e a necessidade de adquirir novas competências, pode também demonstrar ceticismo ou resistência à integração de novas tecnologias baseadas em IA [45]. Superar essa dupla barreira de resistência exige mais do que simples argumentos econômicos sobre ganhos de produtividade futuros. Requer uma comunicação transparente, um forte apelo aos princípios éticos que fundamentam a HCAI – como dignidade humana, autonomia e justiça – e uma visão compartilhada e convincente dos benefícios de longo prazo de um modelo de desenvolvimento tecnológico que seja genuinamente inclusivo e sustentável, alinhado aos princípios da Indústria 5.0 [19].

Nesse contexto, é imperativo que a discussão sobre o futuro do trabalho na era da IA evite as armadilhas de uma polarização simplista entre utopias tecnológicas de abundância sem esforço e distopias de desemprego em massa e controle algorítmico. A abordagem HCAI, fundamentada

na ética e focada na colaboração, oferece um caminho intermediário e pragmático. Ela reconhece o imenso potencial transformador da IA, mas insiste firmemente que seu desenvolvimento e implementação sejam guiados por valores humanos fundamentais e por um compromisso inabalável com a responsabilidade social e o bem-estar coletivo [31]. Portanto, a escolha que se apresenta não é meramente técnica ou econômica; é, em sua essência, uma escolha ética sobre o tipo de futuro que desejamos construir, um futuro onde a tecnologia sirva ao progresso humano em sua acepção mais ampla.

Nesse contexto, as diretrizes discutidas ao longo deste trabalho podem ser sintetizadas de forma prática e aplicada na Tabela II, que organiza princípios centrais da HCAI e exemplos concretos de sua implementação em ambientes sociotécnicos reais.

TABELA II
DIRETRIZES PARA ADOÇÃO ÉTICA DA IA CENTRADA NO HUMANO

Diretriz	Exemplo Prático
Implementar sistemas de IA com supervisão humana significativa	Cobots com painéis de decisão assistida pelo operador no chão de fábrica
Priorizar IA de aumento em vez de substituição	Plataformas de e-learning com IA adaptativa para requalificação interna
Realizar avaliação prévia de impacto no emprego	Análise de risco social obrigatória antes da automação de setores críticos
Fomentar parcerias entre governo, indústria e academia	Programas conjuntos de capacitação para trabalhadores deslocados

VIII. CONCLUSÃO: RUMO A UM FUTURO DE TRABALHO COCRIADO COM A IA

A IA está, inegavelmente, redefinindo o mundo do trabalho. Ignorar as profundas implicações éticas do deslocamento de empregos e da obsolescência de habilidades – particularmente a notória negligência em relação ao suporte necessário aos trabalhadores afetados – seria uma falha moral e social de grandes proporções. Conforme este artigo argumentou, uma resposta ética e sustentável a essa transformação exige uma mudança fundamental de paradigma: da busca, por vezes acrítica, pela automação total para a adoção consciente e deliberada da HCAI.

A HCAI, focada em aumentar as capacidades humanas, promover a colaboração e capacitar os indivíduos, oferece uma visão onde a tecnologia atua como parceira do trabalho humano, e não meramente como sua substituta. Demonstrou-se que essa abordagem, embora tecnicamente viável através de ferramentas de aumento, sistemas colaborativos e IA para capacitação, não pode florescer isoladamente. Ela depende criticamente de um ecossistema de apoio que reconheça sua natureza socio-técnica, envolvendo políticas públicas proativas para incentivar sua adoção e amparar a transição dos trabalhadores, uma reforma educacional orientada para as competências do futuro, e um compromisso genuíno das organizações com a responsabilidade social corporativa.

Construir esse futuro de trabalho cocriado com a IA não será uma tarefa simples. Exigirá diálogo contínuo, colaboração multidisciplinar efetiva entre tecnólogos, cientistas sociais, formuladores de políticas, educadores, empresas

e a sociedade civil. Acima de tudo, demandará uma decisão coletiva e sustentada de priorizar o bem-estar humano, a dignidade no trabalho e a justiça social na trajetória da inovação tecnológica. O desafio é imenso, mas a alternativa – uma sociedade potencialmente mais desigual, fragmentada e instável – torna a busca ativa por este futuro colaborativo não apenas necessária, mas um imperativo ético incontornável.

Este trabalho espera contribuir para a construção de modelos orientadores da prática tecnológica ética, oferecendo subsídios concretos para a formulação de políticas públicas e estratégias organizacionais baseadas no paradigma da IA Centrada no Humano.

IX. LIMITAÇÕES DA PESQUISA E PERSPECTIVAS FUTURAS

Este estudo apresenta uma análise conceitual e ética sobre os impactos da IA no trabalho, com foco na abordagem da HCAI. Entretanto, como um trabalho de natureza teórica, existem limitações a serem reconhecidas.

A principal limitação diz respeito à ausência de validação empírica. Não foram realizados estudos de caso, experimentos controlados nem análises estatísticas que comprovem a aplicabilidade direta das diretrizes propostas. Pesquisas futuras poderão explorar a implementação prática da HCAI em ambientes industriais, plataformas educacionais ou instituições públicas, mensurando indicadores como bem-estar dos trabalhadores, produtividade e aceitação tecnológica.

Outra limitação refere-se ao escopo abrangente da abordagem. Ao integrar conceitos como ética algorítmica, requalificação profissional e Indústria 5.0, optou-se por uma análise holística, o que limita o aprofundamento técnico em cada dimensão. Estudos futuros poderão desenvolver *frameworks* específicos de avaliação do grau de humanocentrismo em sistemas inteligentes, bem como propor indicadores operacionais para orientar políticas públicas e decisões corporativas.

Ademais, este artigo não aborda em profundidade as desigualdades interseccionais relacionadas ao impacto da automação, como as relativas a gênero, raça, faixa etária e localização geográfica. Recomenda-se que investigações futuras incorporem essa dimensão, buscando compreender e mitigar os efeitos assimétricos da transformação digital sobre diferentes grupos sociais.

Por fim, dada a aceleração no desenvolvimento de IA generativa e robótica autônoma, destaca-se a importância de estudos contínuos e atualizados, com foco na governança adaptativa da IA e nos seus efeitos sobre o futuro do trabalho, a justiça social e a inclusão digital.

X. AGRADECIMENTOS

Este trabalho foi apoiado pelas seguintes agências de pesquisa brasileiras: Fundação de Amparo à Pesquisa do Estado de Minas Gerais (Fapemig), Fundação Coordenação de Aperfeiçoamento de Pessoal de Nível Superior (CAPES) e Conselho Nacional de Desenvolvimento Científico e Tecnológico (CNPq).

REFERENCES

- [1] M. Lane, A. Saint-Martin, "The impact of Artificial Intelligence on the labour market," OECD Social, vol. 256, Jan. 2021, doi: 10.1787/7c895724-en.
- [2] D. Marguerit, "Augmenting or automating labor? The effect of AI development on New Work, employment, and wages," Mar. 24, 2025, arXiv: arXiv:2503.19159. doi: 10.48550/arXiv.2503.19159.
- [3] V. Capraro et al., "The impact of generative Artificial Intelligence on socioeconomic inequalities and policy making," PNAS Nexus, vol. 3, no. 6, p. pgae191, Jun. 2024, doi: 10.1093/pnasnexus/pgae191.
- [4] C. B. Frey, M. A. Osborne, "The future of employment: How susceptible are jobs to computerisation?," Technological Forecasting and Social Change, vol. 114, pp. 254–280, Jan. 2017, doi: 10.1016/j.techfore.2016.08.019.
- [5] M. Arntz, T. Gregory, U. Zierahn, "The risk of automation for jobs in OECD countries," OECD Social, vol. 189, May 2016, doi: 10.1787/5jlz9h56dvq7-en.
- [6] E. Brynjolfsson, A. McAfee, "The second machine age: Work, progress, and prosperity in a time of brilliant technologies," Reprint. W. W. Norton And Company, 2016.
- [7] M. Lane, A. Saint-Martin, "The impact of Artificial Intelligence on the labour market: What do we know so far?," OECD Social, Employment and Migration Working Papers, No. 256, 2021, OECD Publishing, Paris, doi: 10.1787/7c895724-en.
- [8] M. Lane, M. Williams, and S. Boeckx, "The impact of AI on the workplace: Main findings from the OECD AI surveys of employers and workers," OECD Social, Employment and Migration Working Papers, vol. 288, Mar. 2023, doi: 10.1787/ea0a0fe1-en.
- [9] A. Georgieff, "Artificial Intelligence and wage inequality," OECD Artificial Intelligence Papers, vol. 13, Apr. 2024, doi: 10.1787/bf98a45c-en.
- [10] J. Schwartz et al., "Ethics and the future of work," Deloitte Insights, 2020. Accessed: Apr. 27, 2025. [Online]. Available: <https://www2.deloitte.com/us/en/insights/focus/human-capital-trends/2020/ethical-implications-of-ai.html>.
- [11] V. Romei, "IMF warns of 'profound concerns' over rising inequality from AI," Financial Times, Jun. 17, 2024. Accessed: Apr. 27, 2025. [Online]. Available: <https://www.ft.com/content/b238e630-93df-4a0c-80d0-fbdf2f13658f>.
- [12] S. Grabowska, S. Saniuk, and B. Gajdzik, "Industry 5.0: Improving humanization and sustainability of Industry 4.0," Scientometrics, vol. 127, no. 6, pp. 3117–3144, Jun. 2022, doi: 10.1007/s11192-022-04370-1.
- [13] J. Alves, T. M. Lima, and P. D. Gaspar, "Is Industry 5.0 a Human-Centred Approach? A Systematic Review," Processes, vol. 11, no. 1, Art. no. 1, Jan. 2023, doi: 10.3390/pr11010193.
- [14] J. M. Rožanec et al., "Human-centric Artificial Intelligence architecture for industry 5.0 applications," International Journal of Production Research, vol. 61, no. 20, pp. 6847–6872, Oct. 2023, doi: 10.1080/00207543.2022.2138611.
- [15] C. Zhao, W. Xu, "Human-AI interaction design standards," Mar. 24, 2025, arXiv: arXiv:2503.16472. doi: 10.48550/arXiv.2503.16472.
- [16] W. Xu, M. J. Dainoff, L. Ge, and Z. Gao, "Transitioning to human interaction with AI systems: New challenges and opportunities for HCI professionals to enable human-centered AI," Mar. 06, 2023, arXiv: arXiv:2105.05424. doi: 10.48550/arXiv.2105.05424.
- [17] X. Xu, Y. Lu, B. Vogel-Heuser, and L. Wang, "Industry 4.0 and Industry 5.0 - Inception, conception and perception," Journal of Manufacturing Systems, vol. 61, pp. 530–535, Oct. 2021, doi: 10.1016/j.jmsy.2021.10.006.
- [18] S. Nahavandi, "Industry 5.0 - A Human-Centric solution," Sustainability, vol. 11, no. 16, Art. no. 16, Jan. 2019, doi: 10.3390/su11164371.
- [19] M. Breque, L. De Nul, A. Petridis, "Industry 5.0: Towards a sustainable, human-centric and resilient European industry," Luxembourg, LU: European Commission, Directorate-General for Research and Innovation, 2021. doi: 10.2777/308407.
- [20] P. K. R. Maddikunta et al., "Industry 5.0: A survey on enabling technologies and potential applications," Journal of Industrial Information Integration, vol. 26, p. 100257, Mar. 2022, doi: 10.1016/j.jii.2021.100257.
- [21] F. Longo, A. Padovano, and S. Umbrello, "Value-oriented and ethical technology engineering in Industry 5.0: A Human-Centric perspective for the design of the factory of the future," Applied Sciences, vol. 10, no. 12, Art. no. 12, Jan. 2020, doi: 10.3390/app10124182.
- [22] S. Huang, B. Wang, X. Li, P. Zheng, D. Mourtzis, and L. Wang, "Industry 5.0 and society 5.0 - Comparison, complementation and co-evolution," Journal of Manufacturing Systems, vol. 64, pp. 424–428, Jul. 2022, doi: 10.1016/j.jmsy.2022.07.010.

- [23] D. Ivanov, "The Industry 5.0 framework: Viability-based integration of the resilience, sustainability, and human-centricity perspectives," *International Journal of Production Research*, Mar. 2023, Accessed: Apr. 27, 2025. [Online]. Available: <https://www.tandfonline.com/doi/abs/10.1080/00207543.2022.2118892>.
- [24] E. Mosqueira-Rey, E. Hernández-Pereira, D. Alonso-Ríos, J. Bobes-Bascarán, and Á. Fernández-Leal, "Human-in-the-loop machine learning: a state of the art," *Artif Intell Rev*, vol. 56, no. 4, pp. 3005–3054, Apr. 2023, doi: 10.1007/s10462-022-10246-w.
- [25] X. Wu, L. Xiao, Y. Sun, J. Zhang, T. Ma, L. He, "A survey of human-in-the-loop for machine learning," *Future Generation Computer Systems*, vol. 135, pp. 364–381, Oct. 2022, doi: 10.1016/j.future.2022.05.014.
- [26] J. Jakubik, D. Weber, P. Hemmer, M. Vössing, and G. Satzger, "Improving the efficiency of human-in-the-Loop systems: Adding artificial to human experts," in *Solutions and Technologies for Responsible Digitalization*, D. Beverungen, C. Lehrer, and M. Trier, Eds., Cham: Springer Nature Switzerland, 2025, pp. 131–147. doi: 10.1007/978-3-031-80122-8_9.
- [27] Z. Huang, Z. Sheng, C. Ma, and S. Chen, "HAIM-DRL: Enhanced human-in-the-loop reinforcement learning for safe and efficient autonomous driving," *Communications in Transportation Research*, vol. 4, p. 100127, Dec. 2024, doi: 10.1016/j.commtr.2024.100127.
- [28] B. Shneiderman, "Human-Centered AI," 1st ed. Oxford University Press/Oxford, 2022. doi: 10.1093/oso/9780192845290.001.0001.
- [29] B. Shneiderman, "Human-Centered Artificial Intelligence: Reliable, safe & trustworthy," *International Journal of Human-Computer Interaction*, vol. 36, no. 6, pp. 495–504, Apr. 2020, doi: 10.1080/10447318.2020.1741118.
- [30] A. Mei, M. Saxon, S. Chang, Z. C. Lipton, and W. Y. Wang, "Users are the north star for AI transparency," Mar. 09, 2023, arXiv: arXiv:2303.05500. doi: 10.48550/arXiv.2303.05500.
- [31] L. Floridi et al., "AI4People — An ethical framework for a good AI society: Opportunities, risks, principles, and recommendations," *Minds & Machines*, vol. 28, no. 4, pp. 689–707, Dec. 2018, doi: 10.1007/s11023-018-9482-5.
- [32] C. Trivedi et al., "Explainable AI for Industry 5.0: Vision, architecture, and potential directions," *IEEE Open J. Ind. Applicat.*, vol. 5, pp. 177–208, 2024, doi: 10.1109/OJIA.2024.3399057.
- [33] S. Schmager, Pappas, Ilias O., and P. and Vassilakopoulou, "Understanding Human-Centred AI: A review of its defining elements and a research agenda," *Behaviour & Information Technology*, pp. 1–40, 2025, doi: 10.1080/0144929X.2024.2448719.
- [34] J. C. S. Lima, "Ethical challenges in Artificial Intelligence: Navigating the boundaries of responsible innovation," unpublished.
- [35] T. Islam, S. Afrin, "Mitigating AI bias in recruitment: Policy approaches for transparent candidate selection and broader implications for trust in algorithmic decisions," 2023, [Online]. Accessed: Apr. 27, 2025. Available: <https://www.esrpjta.org/jeta-v3i7p115>.
- [36] L. Floridi, M. Taddeo, "What is data ethics?," *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences*, vol. 374, no. 2083, p. 20160360, Dec. 2016, doi: 10.1098/rsta.2016.0360.
- [37] B. D. Mittelstadt, P. Allo, M. Taddeo, S. Wachter, and L. Floridi, "The ethics of algorithms: Mapping the debate," *Big Data & Society*, vol. 3, no. 2, p. 2053951716679679, Dec. 2016, doi: 10.1177/2053951716679679.
- [38] M. Jovanović, "The ethical implications of Artificial Intelligence in the workplace: Bias, discrimination, and transparency," *International Conference "Actual economy: Local solutions for global challenges"*, pp. 188–190, Nov. 2024, [Online]. Accessed: Apr. 29, 2025. Available: <https://conferances.com/index.php/journal/article/view/328>.
- [39] D. Acemoglu, P. Restrepo, "Automation and new tasks: How technology displaces and reinstates labor," *Journal of Economic Perspectives*, vol. 33, no. 2, pp. 3–30, May 2019, doi: 10.1257/jep.33.2.3.
- [40] J. Danaher, "The rise of the robots and the crisis of moral patiency," *AI Soc.*, vol. 34, no. 1, pp. 129–136, Mar. 2019, doi: 10.1007/s00146-017-0773-9.
- [41] B. C. Stahl, "Responsible innovation ecosystems: Ethical implications of the application of the ecosystem concept to artificial intelligence," *International Journal of Information Management*, vol. 62, p. 102441, Feb. 2022, doi: 10.1016/j.ijinfomgt.2021.102441.
- [42] D. H. Autor, "Why are there still so many jobs? The history and future of workplace automation," *Journal of Economic Perspectives*, vol. 29, no. 3, pp. 3–30, Sep. 2015, doi: 10.1257/jep.29.3.3.
- [43] B. Moll, L. Rachel, and P. Restrepo, "Uneven growth: Automation's impact on income and wealth inequality," *Econometrica*, vol. 90, no. 6, pp. 2645–2683, 2022, doi: 10.3982/ECTA19417.
- [44] K. C. Kellogg, M. A. Valentine, and A. Christin, "Algorithms at work: The new contested terrain of control," *ANNALS*, vol. 14, no. 1, pp. 366–410, Jan. 2020, doi: 10.5465/annals.2018.0174.
- [45] U. Leicht-Deobald et al., "The challenges of algorithm-based HR decision-making for personal integrity," *J Bus Ethics*, vol. 160, no. 2, pp. 377–392, Dec. 2019, doi: 10.1007/s10551-019-04204-w.
- [46] F. Fossen and A. Sorgner, "Mapping the future of occupations: Transformative and destructive effects of new digital technologies on jobs," *Foresight and STI Governance (Foresight-Russia till No. 3/2015)*, vol. 13, no. 2, pp. 10–18, 2019, [Online]. Accessed: Apr. 28, 2025. Available: <https://www.researchgate.net/publication/334634967>.
- [47] D. A. Spencer, "Fear and hope in an age of mass automation: debating the future of work," *New Technology, Work and Employment*, vol. 33, no. 1, Art. no. 1, Mar. 2018, [Online]. Accessed: Apr. 28, 2025. Available: <http://eprints.whiterose.ac.uk/126551/>.
- [48] E. Malone, S. Afroogh, J. DCruz, and K. R. Varshney, "When trust is zero sum: Automation threat to epistemic agency," Aug. 19, 2024, arXiv: arXiv:2408.08846. doi: 10.48550/arXiv.2408.08846.
- [49] D. Acemoglu and P. Restrepo, "Tasks, automation, and the rise in U.S. wage inequality," *Econometrica*, vol. 90, no. 5, pp. 1973–2016, 2022, doi: 10.3982/ECTA19815.
- [50] B. Shneiderman, "Bridging the gap between ethics and practice: Guidelines for reliable, safe, and trustworthy Human-centered AI systems," *ACM Trans. Interact. Intell. Syst.*, vol. 10, no. 4, p. 26:1–26:31, Outubro 2020, doi: 10.1145/3419764.
- [51] M. Passalacqua et al., "Human-centred AI in Industry 5.0: a systematic review," *International Journal of Production Research*, vol. 63, no. 7, pp. 2638–2669, Apr. 2025, doi: 10.1080/00207543.2024.2406021.
- [52] Y. K. Dwivedi et al., "Opinion paper: 'So what if ChatGPT wrote it?' Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy," *International Journal of Information Management*, vol. 71, p. 102642, Aug. 2023, doi: 10.1016/j.ijinfomgt.2023.102642.
- [53] J. Fiegler-Rudol, K. Lau, A. Mroczek, and J. Kasperczyk, "Exploring Human-AI dynamics in enhancing workplace health and safety: A narrative review," *International Journal of Environmental Research and Public Health*, vol. 22, no. 2, Art. no. 2, Feb. 2025, doi: 10.3390/ijerph22020199.
- [54] L. Waardenburg, "Human-AI collaboration: A blessing or a curse for safety at work?," *Tecnoscienza – Italian Journal of Science & Technology Studies*, vol. 15, no. 1, Art. no. 1, Jul. 2024, doi: 10.6092/issn.2038-3460/19964.
- [55] M. C. Elish, "Moral crumple zones: Cautionary tales in human-robot interaction (pre-print)," Mar. 01, 2019, Social Science Research Network, Rochester, NY: 2757236. doi: 10.2139/ssrn.2757236.
- [56] F. G. Antonaci et al., "Workplace well-being in Industry 5.0: A worker-centered systematic review," *Sensors*, vol. 24, no. 17, Art. no. 17, Jan. 2024, doi: 10.3390/s24175473.
- [57] J. Manyika et al., "Jobs of the future: Jobs lost, jobs gained," McKinsey & Company, 2017. Accessed: Apr. 28, 2025. [Online]. Available: <https://www.mckinsey.com/featured-insights/future-of-work/jobs-lost-jobs-gained-what-the-future-of-work-will-mean-for-jobs-skills-and-wages>.
- [58] V. Dignum, "Taking responsibility," in *Responsible Artificial Intelligence: How to Develop and Use AI in a Responsible Way*, V. Dignum, Ed., Cham: Springer International Publishing, 2019, pp. 47–69. doi: 10.1007/978-3-030-30371-6_4.
- [59] OECD, "AI and the future of skills, Volume 1: Capabilities and assessments," *Educational Research and Innovation*, OECD Publishing, Paris, 2021, doi: 10.1787/5ee71f34-en.
- [60] WEF, "The future of jobs report 2023," *World Economic Forum*, May 2023. Accessed: Apr. 28, 2025. [Online]. Available: <https://www.weforum.org/publications/the-future-of-jobs-report-2023/>.
- [61] D. Acemoglu, P. Restrepo, "The wrong kind of AI? Artificial Intelligence and the future of labor demand," Mar. 2019, National Bureau of Economic Research: 25682. doi: 10.3386/w25682.
- [62] E. Brynjolfsson, "The turing trap: The promise & peril of human-like Artificial Intelligence," *Daedalus*, vol. 151, no. 2, pp. 272–287, May 2022, doi: 10.1162/daed_a.01915.
- [63] D. H. Autor, "The labor market impacts of technological change: From unbridled enthusiasm to qualified optimism to vast uncertainty," May 01, 2022, Social Science Research Network, Rochester, NY: 4122803. Accessed: Apr. 28, 2025. [Online]. Available: <https://papers.ssrn.com/abstract=4122803>.
- [64] J. Galvin and L. LaBerge, "Digital strategy in the postpandemic era," May, 2021, McKinsey Global Institute, Accessed: Apr. 28, 2025. [Online]. Available: <https://www.mckinsey.com/capabilities/mckinsey-digital/our-insights/the-new-digital-edge-rethinking-strategy-for-the-postpandemic-era>.

ARTIGO III: OTIMIZAÇÃO SOCIO-TÉCNICA DE TERMINAIS INTERMODAIS: UMA ABORDAGEM DE SIMULAÇÃO MULTI-OBJETIVO PARA EQUILIBRAR EFICIÊNCIA, BEM-ESTAR DO TRABALHADOR E EQUIDADE

Otimização Socio-Técnica de Terminais Intermodais: Uma Abordagem de Simulação Multi-Objetivo para Equilibrar Eficiência, Bem-Estar do Trabalhador e Equidade

1st Julio C. S. Lima
Programa de Pós-Graduação em
Engenharia de Sistemas e Automação
Universidade Federal de Lavras
Lavras/MG, Brasil
lima.julio.cs@gmail.com

2nd Danton Diego Ferreira
Departamento de Automática
Universidade Federal de Lavras
Lavras/MG, Brasil
danton@ufla.br

3rd Felipe O. Silva
Departamento de Automática
Universidade Federal de Lavras
Lavras/MG, Brasil
felipe.oliveira@ufla.br

Resumo—A otimização de terminais intermodais tem se concentrado historicamente em métricas de eficiência técnica, como maximização da vazão e minimização de tempos de ciclo, negligenciando os impactos socio-técnicos sobre os trabalhadores e *stakeholders* externos. Esta omissão configura uma “lacuna de tradução de valores”, onde princípios éticos como bem-estar e justiça não são operacionalizados como objetivos de otimização. Para endereçar esta lacuna, este artigo apresenta uma metodologia para a otimização multi-objetivo de um terminal intermodal de grãos que integra explicitamente considerações éticas. Utilizando simulação de eventos discretos na plataforma JaamSim, desenvolvemos um modelo que, além das métricas operacionais tradicionais, quantifica (1) a fadiga do operador, através de um submodelo que abstrai o processo homeostático de acumulação e recuperação de cansaço, e (2) a equidade na distribuição dos tempos de espera dos caminhoneiros, medida pelo Coeficiente de Gini. Acoplamos este modelo socio-técnico ao algoritmo genético NSGA-II para explorar a fronteira de Pareto e mapear os *trade-offs* quantitativos entre cinco objetivos concorrentes: tempo de permanência, vazão, utilização de ativos, fadiga média dos operadores e desigualdade nos tempos de espera. Os resultados demonstram a viabilidade de se modelar e otimizar para estes objetivos conflitantes, fornecendo aos gestores um “menu de políticas” que torna as consequências éticas das decisões operacionais explícitas e mensuráveis. A análise da fronteira de Pareto revela que ganhos em produtividade estão frequentemente associados a um aumento na fadiga e na desigualdade, quantificando o custo humano da eficiência. Este trabalho contribui com uma metodologia prática para o projeto e gestão de sistemas logísticos mais justos e humanocêntricos, alinhados aos princípios da Indústria 5.0, ao transformar valores éticos em objetivos de otimização quantificáveis.

Index Terms—Simulação de Eventos Discretos, Otimização Multi-Objetivo, NSGA-II, Logística Humanocêntrica, Indústria 5.0, Fadiga do Operador, Justiça em Filas, Coeficiente de Gini.

I. INTRODUÇÃO

Terminais intermodais representam elos vitais e complexos nas cadeias de suprimentos globais, particularmente no escoamento de commodities como grãos, onde a eficiência operacional e a minimização de custos têm sido, historicamente, os vetores primordiais de decisão [1]. A otimização desses sistemas, frequentemente abordada por meio de simulação de eventos discretos (DES), tem se concentrado em métricas

de desempenho essencialmente técnicas, como a redução de tempos de ciclo, a maximização da utilização de recursos e a minimização do comprimento de filas [2], [3]. Essa abordagem, embora crucial para a competitividade econômica, revela-se cada vez mais limitada em um cenário contemporâneo que demanda maior responsabilidade socio-técnica.

A busca incessante pela eficiência, quando desvinculada de considerações éticas, pode ocorrer em detrimento dos impactos sobre os *stakeholders* humanos do sistema. Questões fundamentais como a justiça na alocação de tarefas e no acesso a recursos compartilhados — por exemplo, a priorização no atendimento a caminhões — permanecem frequentemente subexploradas em modelos de otimização [4], [5]. Similarmente, o bem-estar dos trabalhadores, incluindo motoristas submetidos a longas e incertas esperas, e a transparência dos processos decisórios, especialmente quando mediados por algoritmos, são dimensões éticas que raramente são integradas como variáveis centrais na avaliação de desempenho desses terminais [6]–[11]. Esta lacuna configura-se não apenas como uma omissão analítica, mas como um desafio ético premente em uma era de transição para a Indústria 4.0 e emergência da Indústria 5.0, que preconizam uma maior centralidade humana e a integração de valores societários nos sistemas produtivos e logísticos [11]–[17].

Nesse contexto, torna-se essencial ancorar a discussão de justiça operacional em uma base conceitual sólida. A literatura clássica de pesquisa operacional em sistemas de filas oferece contribuições fundamentais para entender como se estrutura o *trade-off* entre eficiência e equidade no atendimento. Estudos como os de Avi-Itzhak e Levy [18] e Raz et al. [19] introduziram o *Resource Allocation Queueing Fairness Measure* (RAQFM), uma métrica que considera simultaneamente o tempo de espera (senioridade) e o tempo de serviço das tarefas, oferecendo critérios formais para avaliar a justiça sob diferentes disciplinas de fila, como “Primeiro a Chegar, Primeiro a Ser Servido” (*First-Come, First-Served* - FCFS), “Compartilhamento de Processador” (*Processor Sharing* - PS), “Último a Chegar, Primeiro a Ser Servido” (*Last-Come, First-Served* - LCFS) e “Ordem Aleatória de Serviço” (*Ran-*

dom Order of Service - ROS). Em complemento, Wierman e Harchol-Balter [20] mostraram como o *slowdown* — razão entre tempo de resposta e tamanho da tarefa — quantifica o impacto da variabilidade na priorização, destacando o conflito estrutural entre maximização do tempo de processamento e justiça percebida. Ainda, Whitt [21] investigou o fenômeno de *overtaking*, evidenciando como políticas não-FCFS podem exacerbar desigualdades de espera. Esses fundamentos foram estendidos para cenários de múltiplas filas [22] e para redes de comunicação, onde Kelly et al. [23] e Bonald et al. [24] conectaram a justiça proporcional ao balanceamento de recursos em regimes multi-classe.

Assim, além de discussões normativas, há uma tradição matemática consolidada para a quantificação da justiça em filas, que respalda a adoção de métricas robustas, como o Coeficiente de Gini dos tempos de espera. Essa conexão entre teoria de filas e métricas de desigualdade fortalece a justificativa metodológica deste estudo, que propõe ir além do tradicional “Primeiro a Chegar, Primeiro a Ser Servido” (*First-In, First-Out* - FIFO) e explicitar o *trade-off* entre eficiência operacional e justiça distributiva no ambiente do terminal.

Para endereçar essa lacuna, este artigo propõe uma abordagem metodológica para o desenvolvimento de um modelo de simulação de eventos discretos (DES) eticamente integrado. O foco recai sobre a modelagem de um terminal intermodal de grãos, um sistema caracterizado por fluxos complexos e recursos compartilhados [2], [25], [26]. O trabalho avança para além da otimização tradicional ao incorporar, como objetivos explícitos, métricas quantificáveis para a **fadiga do operador**, representando o bem-estar do trabalhador e a segurança operacional, e a **equidade na distribuição dos tempos de espera**, representando a justiça distributiva para com os *stakeholders* externos (caminhoneiros). Utilizando a plataforma open-source JaamSim [27]–[30] e o algoritmo multi-objetivo NSGA-II [31], a pesquisa busca mapear de forma transparente os *trade-offs* entre desempenho técnico e impacto humano.

As principais contribuições deste estudo são multifacetadas. Primeiramente, propõe-se uma metodologia inovadora para a operacionalização de princípios éticos em modelos DES aplicados a sistemas logísticos. Em segundo lugar, o trabalho consolida métricas formais para justiça em filas, alinhadas a medidas como o Coeficiente de Gini, e fundamenta o submodelo de fadiga em bases de ergonomia cognitiva. Por fim, a análise quantitativa dos *trade-offs* visa instrumentalizar gestores e formuladores de políticas com uma compreensão mais completa das implicações socio-técnicas de suas decisões, fomentando práticas alinhadas aos ideais de uma Indústria 5.0 mais justa, centrada no ser humano.

II. O SISTEMA DO TERMINAL INTERMODAL: PROCESSOS E RECURSOS

Para desenvolver um modelo de simulação robusto e representativo, é imperativo, primeiramente, decompor o sistema real em seus componentes fundamentais: os processos, os recursos e as lógicas de decisão que governam o fluxo de entidades. Esta seção estabelece essa base factual para o terminal intermodal de grãos, mapeando a realidade operacional

para os construtos do modelo de simulação desenvolvido na plataforma JaamSim. O objetivo é garantir a transparência e a validade do modelo, tornando-o uma base defensável para a subsequente análise de otimização.

A. Descrição Geral do Processo

A operação de um terminal intermodal de grãos envolve a gestão precisa do fluxo de veículos e da movimentação de cargas para garantir que todas as etapas do processo ocorram de forma sincronizada e produtiva. O macrofluxo de um caminhão para descarga de grãos é tipicamente estruturado em uma sequência de processos e checkpoints bem definidos, que se inicia com a chegada do veículo e termina com sua saída, após a entrega da carga. O percurso pode ser resumido nas seguintes etapas principais:

- **Portaria de Entrada:** Controle de acesso e registro inicial dos veículos (Fig. 1a);
- **Classificação:** Etapa essencial onde a qualidade dos grãos é determinada por meio de amostragem (calagem) e análise laboratorial. O resultado desta etapa influencia a destinação da carga (Fig. 1b);
- **Primeira Pesagem (Peso Bruto):** O caminhão carregado é pesado para registrar seu peso total (Fig. 1c).
- **Descarga:** Realizada em equipamentos especializados, como os tombadores, que esvaziam o conteúdo do caminhão de forma rápida e eficiente (Fig. 1d);
- **Segunda Pesagem (Tara):** O caminhão, agora vazio, é pesado novamente para aferir a tara e, por diferença, determinar o peso líquido do produto entregue (Fig. 1c);
- **Portaria de Saída:** Registro da saída e liberação final do veículo (Fig. 1a).

Cada uma dessas etapas é realizada em locais físicos distintos dentro do terminal, utilizando equipamentos específicos e mão de obra qualificada.

B. Mapeamento do Fluxo de Processo

A construção do modelo de simulação no ambiente JaamSim [32] parte da decomposição das atividades do terminal intermodal em processos discretos, organizados em sequência lógica e operacional. A Fig. 2 apresenta o macrofluxo do processo de descarga rodoviária de grãos, tal como representado no modelo.

Neste fluxograma, observa-se o trajeto percorrido por um caminhão carregado desde o acesso à portaria de entrada até a saída do terminal, após o descarregamento completo da carga. O processo envolve as seguintes etapas: (i) entrada no terminal, (ii) classificação da carga, (iii) pesagem inicial (peso bruto), (iv) descarga por tombamento, (v) pesagem final (peso da tara) e (vi) saída do terminal. Após a descarga, o produto é transferido para o sistema de armazenagem, de onde será posteriormente encaminhado ao modal ferroviário para exportação.

Embora o modelo computacional represente esses processos com maior granularidade — incluindo filas intermediárias, alocação de recursos e lógica condicional baseada em disponibilidade —, opta-se por apresentar apenas o macrofluxo nesta seção. As especificidades do modelo, como o uso de objetos *Queue*, *Server*, *ResourcePool* e *Branch*, serão detalhadas ao longo do artigo, conforme os módulos e experimentos forem discutidos.



(a) Portaria Entrada/Saída



(b) Classificação - Equipamento Calagem e Deslonamento



(c) Balança Rodoviária - Peso Bruto/Tara



(d) Tombador - Descarga Rodoviária

Fig. 1: Processo de Descarga de Grãos

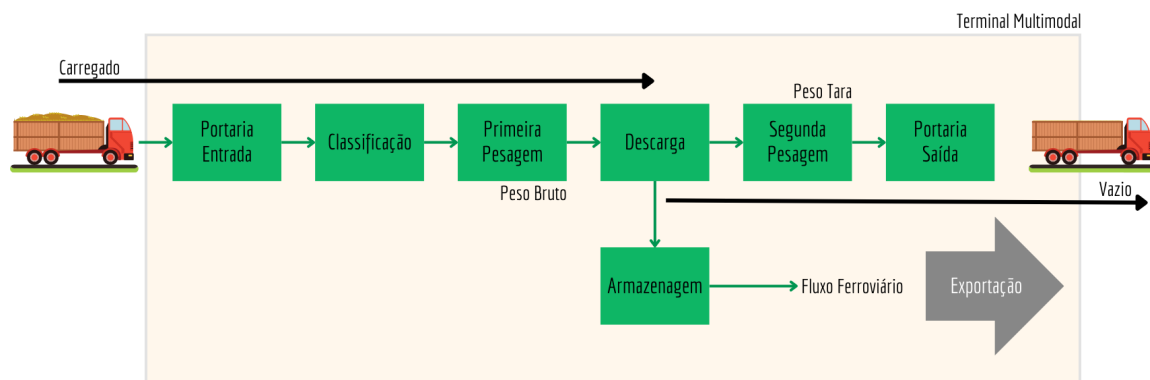


Fig. 2: Macro Fluxo Rodoviário

Este mapeamento de alto nível fornece o contexto essencial para a compreensão das decisões de modelagem e da lógica operacional que será analisada e otimizada nas seções seguintes.

C. Definição de Recursos e Entidades

O sistema é composto por entidades que fluem através dele e recursos que são utilizados para processar essas entidades.

- **Entidades:** A entidade primária que se move através do sistema é o Truck, definida no modelo como

`SimEntity { Truck }`. Cada caminhão possui atributos que podem ser modificados ao longo do processo, como o `NetWeigh` (peso líquido).

- **Recursos (Equipamentos):** Os equipamentos fixos do terminal são modelados como objetos `Server` (para processos com tempo de serviço) ou `Resource` (para recursos com capacidade finita que são “tomados” e “liberados”). Exemplos incluem `EntranceGate01`, `EntranceGate02`, `Discharger01`, `Discharger02`,

Discharger03, Discharger04,
ResourceFirstWeighingScale e
ResourceSecondWeighingScale.

- **Recursos (Humanos):** Uma característica distintiva deste modelo é a representação explícita da força de trabalho. Os operadores não são recursos passivos, mas são modelados como `ResourceUnits` individuais, agrupados em `ResourcePools`. Por exemplo, `ResourcePoolSampler` agrupa os amostradores, `ResourcePoolClassificationStation01` agrupa os classificadores da estação 1, e `ResourcePoolDischarger01` agrupa os operadores do tombador 1. Crucialmente, cada `ResourceUnit` (p. ex., `classifier_1_T1`, `OperatorDischarge01_T2`) possui atributos dinâmicos, como `NivelFadiga`, `TaxaAcumulo` e `TaxaRecuperacao`, e está associado a um turno de trabalho específico (`ShiftThreshold1`, `ShiftThreshold2`, `ShiftThreshold3`). Essa modelagem transforma a simulação de um sistema puramente mecânico em uma simulação socio-técnica, permitindo a análise do impacto da alocação de pessoal e da fadiga no desempenho geral.

Para garantir total transparência e auditabilidade, a Tabela I a seguir estabelece uma correspondência direta entre as etapas do processo do mundo real, os recursos envolvidos e os objetos específicos nomeados no código do modelo JaamSim.

D. Limites do Sistema e Externalidades

A definição clara dos limites do sistema é fundamental para a interpretação correta dos resultados da simulação [33], [34]. Neste estudo, o sistema inicia com a geração de caminhões no objeto `Generator`, que simula a sua chegada de um pátio de espera externo, e termina quando o caminhão é processado pelo objeto `EntitySink1` após sair pela portaria. Essa delimitação é necessária para criar um modelo computacionalmente tratável.

No entanto, é crucial reconhecer que essa fronteira sistêmica exclui externalidades significativas, que são os efeitos não intencionais das atividades do terminal sobre terceiros não diretamente envolvidos na transação [35], [36]. A principal externalidade negativa neste contexto é o impacto da fila de caminhões que se forma fora dos portões do terminal. Essa fila pode causar congestionamento no tráfego local, poluição do ar e sonora devido aos motores ligados, e transtornos gerais para a comunidade circundante [37]. Embora o modelo não quantifique diretamente esses impactos (p. ex., emissões de CO₂, decibéis de ruído), a variável de entrada.

`ConditionTruckArrivalMaxTrucks` atua como um proxy, representando uma política do terminal para limitar o número de caminhões que podem se aproximar, reconhecendo implicitamente essa externalidade. A discussão sobre esses custos sociais externalizados será retomada na análise ética dos resultados, pois uma política operacional que minimiza os tempos de fila dentro do terminal ao custo de maximizar a fila fora dele pode ser eficiente do ponto de vista da empresa, mas eticamente problemática do ponto de vista da comunidade.

III. ESTRUTURA ÉTICA PARA A ANÁLISE DE OPERAÇÕES DO TERMINAL

A transição de uma análise puramente técnica para uma avaliação socio-técnica exige a construção de um arcabouço teórico que justifique a inclusão de métricas não tradicionais. Esta seção estabelece esse alicerce, demonstrando que os conceitos de “equidade” e “fadiga” não são escolhas arbitrárias, mas sim construtos robustos, fundamentados na literatura de teoria da justiça, fatores humanos e ergonomia cognitiva. O objetivo é legitimar sua utilização como funções-objetivo na otimização, transformando-os de preocupações qualitativas em variáveis quantificáveis e centrais para a análise de desempenho do terminal.

A. Equidade (Fairness) em Sistemas de Fila: Para Além do FIFO

A disciplina de fila mais comum e intuitivamente “justa” é a FIFO [38], [39]. A sua aparente justiça reside na simplicidade e transparência de sua regra. No entanto, uma análise mais profunda, informada pela teoria da justiça organizacional, revela uma distinção crucial entre justiça processual e justiça distributiva [40], [41]. A justiça processual refere-se à percepção de que os procedimentos e políticas utilizados para tomar decisões são justos, enquanto a justiça distributiva se refere à percepção de que os resultados ou alocações finais são justos [42], [43]. O FIFO é um exemplo de alta justiça processual, pois a regra é clara e universalmente aplicada. Contudo, ele pode levar a resultados de justiça distributiva questionáveis.

O principal problema do FIFO em sistemas heterogêneos é que ele ignora um *trade-off* fundamental: senioridade (tempo de espera) vs. demanda de serviço (tempo de processamento) [38]. Em um terminal de grãos, diferentes caminhões podem ter demandas de serviço muito distintas (p. ex., um caminhão padrão vs. um que exige um procedimento de classificação especial e demorado). Sob uma regra FIFO estrita, um caminhão com uma demanda de serviço curta pode ficar retido por um tempo desproporcionalmente longo atrás de um caminhão com uma demanda de serviço longa que chegou apenas alguns instantes antes. Isso não apenas pode aumentar o tempo médio de espera geral, mas, crucialmente, aumenta a variabilidade e a desigualdade dos tempos de espera entre os usuários. Para um caminhoneiro, a previsibilidade do tempo de espera é tão importante quanto sua duração média. Uma alta variabilidade significa incerteza e pode ser percebida como uma distribuição injusta dos “custos” de espera.

Para capturar essa dimensão de justiça distributiva, este estudo propõe ir além da simples medição do tempo médio de espera e adotar métricas que quantifiquem a desigualdade na distribuição desses tempos. Três métricas são candidatas principais: a variância, o Coeficiente de Variação (CV) e o Coeficiente de Gini (Tabela II).

A variância do tempo de espera, definida como a média dos desvios quadrados da média, mede a dispersão dos tempos individuais em torno da média; uma variância maior indica maior desigualdade [44]. Já o Coeficiente de Variação (CV), que é o desvio padrão dividido pela média, padroniza a variância, tornando-a independente da escala, o que facilita

TABELA I: Mapeamento de Processos, Recursos e Construtos do Modelo JaamSim.

Etapa	Processo (Real)	Descrição da Atividade	Recursos (Humanos/Equipamentos)	Chave (Humana)	Objetos JaamSim Correspondentes
1.	Fila Externa	Espera para acesso ao terminal	N/A (Fora do limite do sistema)		Generator, EntityConveyor { ParkingToTerminal }
2.	Portaria de Entrada	Registro e verificação de documentos	Porteiros (implícito), Gates de Entrada		Queue { QueueTerminal }, EntityGate { GateTerminal }, Queue { QueueEntranceGate }, Server { EntranceGate01, EntranceGate02 }
3.	Classificação	Amostragem e análise da qualidade do grão	Amostradores (ResourcePool-Sampler), Classificadores (ResourcePoolClassification-Station*), Calador, Equip. de Análise		Queue { QueueGateClassification }, Branch { RouteClassificationDecision }, Server { Sampling }, Server { ClassificationStation* }
4.	Pesagem Bruta	Pesagem do caminhão carregado	Operador de Balança (implícito), Balança de Entrada		Queue { QueueFirstWeighing }, Seize { SeizeFirstWeighingScale }, EntityDelay { FirstWeighing }, Resource { ResourceFirstWeighingScale }
5.	Descarga	Posicionamento e tombamento do caminhão	Operadores de Tombador (ResourcePoolDischarger*), Tombadores		Queue { QueueGateDischarger* }, Server { Discharger01, Discharger02, Discharger03, Discharger04 }
6.	Pesagem de Tara	Pesagem do caminhão vazio	Operador de Balança (implícito), Balança de Saída		Queue { QueueSecondWeighing }, Seize { SeizeSecondWeighingScale }, EntityDelay { SecondWeighing }, Resource { ResourceSecondWeighingScale }
7.	Portaria de Saída	Liberação do caminhão	Porteiros (implícito), Gates de Saída		Queue { QueueExitGate }, Server { ExitGate01, ExitGate02 }, EntitySink { EntitySink1 }

a comparação entre cenários com diferentes magnitudes de tempo médio. Por fim, o Coeficiente de Gini é uma medida consolidada e amplamente utilizada para quantificar a desigualdade de uma distribuição [45], [46]. Originalmente desenvolvido para medir a desigualdade de renda, o Gini pode ser aplicado a qualquer distribuição de recursos, incluindo o “recurso” de tempo de espera. Um Coeficiente de Gini de 0 representa a igualdade perfeita (todos os caminhoneiros esperam exatamente o mesmo tempo), enquanto um valor próximo de 1 indica a desigualdade máxima (uma pequena fração dos caminhoneiros tem tempos de espera muito curtos, enquanto a grande maioria enfrenta esperas extremamente longas) [46]–[48].

Ao incorporar a minimização de métricas como a variância, o CV ou o Coeficiente de Gini dos tempos de espera como funções-objetivo (por exemplo, f_5), o modelo torna a escolha entre justiça processual (implícita no FIFO) e justiça distributiva (quantificada por essas métricas) explícita, mensurável e passível de otimização.

B. Fadiga do Operador como Fator Humano Crítico e Métrica de Bem-Estar

Operações industriais contínuas (24/7), como as de terminais portuários e de transbordo, impõem demandas significativas aos trabalhadores, criando um ambiente onde a fadiga se torna um fator de risco crítico [49]. Em operações industriais contínuas (24/7), a fadiga transcende o mero cansaço, constituindo-se em um estado fisiológico que comprovadamente degrada o desempenho cognitivo e motor, diminuindo a vigilância, retardando o tempo de reação e aumentando a probabilidade de erro humano [50]. Em um ambiente onde a operação de equipamentos pesados é a norma, um erro induzido pela fadiga pode ter consequências catastróficas [51]. Estudos quantitativos em terminais de contêineres, com funções análogas às do terminal de grãos,

demonstraram que a fadiga simultânea de operadores-chave pode reduzir a eficiência geral do terminal em mais de 17% [52].

Para capturar esse fenômeno de forma cientificamente defensável, o submodelo de fadiga implementado neste estudo, embora uma abstração prática, é fundamentado em modelos biomatemáticos de fadiga (BMMFs) bem estabelecidos [53]–[55]. Esses modelos são ferramentas quantitativas que preveem os níveis de alerta e desempenho com base em fatores como histórico de sono-vigília e ritmos biológicos [56], [57]. A base teórica para muitos BMMFs é o Modelo de Três Processos de Alerta, desenvolvido por Åkerstedt e Folkard [58]. Este modelo postula que o estado de alerta de uma pessoa é determinado pela interação de três componentes:

- **Processo S (Homeostático):** Este processo representa a “pressão do sono”. Funciona como um cronômetro que se acumula durante o período de vigília e se dissipa durante o sono. Quanto mais tempo uma pessoa fica acordada, maior a pressão do sono e, conseqüentemente, menor o seu nível de alerta. A recuperação deste processo é exponencial durante o sono [58], [59].
- **Processo C (Circadiano):** Este é o marcapasso biológico endógeno, o nosso “relógio biológico”. Ele gera uma flutuação rítmica no alerta ao longo de um período de aproximadamente 24 horas, independentemente do débito de sono. Este processo explica por que a sonolência atinge um pico nas primeiras horas da manhã (tipicamente entre 02:00 e 06:00), mesmo que um indivíduo tenha dormido o suficiente no dia anterior. É o principal fator que torna o trabalho noturno tão desafiado [58], [59].
- **Processo W (Inércia do Sono):** Este componente descreve o período transitório de sonolência, desorientação e desempenho cognitivo reduzido que ocorre imediata-

TABELA II: Análise Comparativa de Métricas de Justiça para Tempos de Espera.

Métrica	Definição	Propriedades e Sensibilidades Chave	Prós para Análise de Filas	Contras para Análise de Filas
Variância	Média dos desvios quadrados da média.	Sensível a outliers, dependente da escala.	Simples de calcular e interpretar.	Não é independente da escala; difícil de comparar entre cenários com diferentes tempos médios de espera.
Coefficiente de Variação (CV)	Desvio dividido pela Média.	Independente da escala. Altamente sensível a outliers extremos (cauda direita).	Bom para penalizar cenários com algumas esperas catastróficas longas.	Pode ser instável se a média for próxima de zero. Menos foco na maior parte da distribuição.
Coefficiente de Gini	Baseado na curva de Lorenz; metade da diferença média absoluta relativa.	Independente da escala. Mais sensível a mudanças no meio da distribuição do que nas caudas.	Robusto, padrão amplamente utilizado para desigualdade. Captura bem a desigualdade geral da distribuição.	Menos sensível a alguns casos de sofrimento extremo, que podem ser um aspecto chave da injustiça percebida. Pode ser menos intuitivo que a variância.

mente após o despertar. A inércia do sono se dissipa exponencialmente nos primeiros minutos a horas de vigília [59].

No presente estudo, o submodelo de fadiga implementado no JaamSim é uma abstração simplificada, restringindo-se exclusivamente à dinâmica homeostática da fadiga, com base no denominado Processo S. Assim, os parâmetros *TaxaAcumulo* e *TaxaRecuperacao* é considerada constante ao longo do tempo, independentemente do horário de trabalho. Consequentemente, no modelo, períodos laborais diurnos e noturnos são tratados como igualmente fatigantes, o que omite completamente o componente circadiano (Processo C) que regula flutuações de alerta e fadiga ao longo das 24 horas. Essa decisão de simplificação operacionaliza apenas uma dimensão da fisiologia do sono e da fadiga, diferentemente de modelos mais robustos, como o SAFTE-FAST (*Synthetic Apprentice Fatigue and Task Scheduling - FAST*) [60], amplamente utilizado na aviação para avaliar o impacto da fadiga em operações aéreas, que incorporam explicitamente a interação entre os processos homeostático e circadiano.

A implementação no modelo JaamSim, com seus parâmetros *TaxaAcumulo* e *TaxaRecuperacao* associados a cada tarefa e os horários de turno, funciona como uma abstração computacionalmente eficiente dessa dinâmica. A *TaxaAcumulo* pode ser vista como uma combinação da pressão do Processo S e da carga de trabalho da tarefa, enquanto a *TaxaRecuperacao* durante as pausas e a redefinição da fadiga entre os turnos simulam a dissipação da pressão do sono. Ao citar o Modelo de Três Processos como o fundamento teórico, a abordagem do estudo ganha em rigor científico, movendo-se de uma heurística ad-hoc para uma representação simplificada, mas teoricamente justificada, do fenômeno da fadiga.

Além de ser um fator de risco, a fadiga é um poderoso *proxy* para o bem-estar do trabalhador. Níveis cronicamente elevados de fadiga estão associados a estresse, problemas de saúde e baixa satisfação no trabalho. Portanto, ao incluir a minimização da fadiga média dos operadores como uma função-objetivo (f_4), o estudo alinha-se diretamente aos princípios da IA Centrada no Humano (HCAI) e da Indústria 5.0, que reposicionam o ser humano e seu bem-estar no centro do projeto de sistemas produtivos.

C. Transparência e Responsabilidade Algorítmica

A introdução de um sistema de otimização como o NSGA-II, que recomenda configurações operacionais, eleva a discussão para o campo da responsabilidade algorítmica [6], [11], [16]. Se o algoritmo sugere uma política que, embora eficiente, aumenta drasticamente a fadiga dos operadores ou a desigualdade entre os caminhoneiros, surge a questão da responsabilidade: é do algoritmo, do desenvolvedor do modelo, ou do gerente que aceita a recomendação?

Essa questão é central no campo emergente da gestão algorítmica, onde decisões tradicionalmente humanas sobre alocação de trabalho, agendamento e priorização de serviço são delegadas a sistemas automatizados [11], [61], [62]. Tal prática acarreta riscos legais e éticos significativos, incluindo o potencial para discriminação algorítmica (se o sistema aprende a favorecer certos tipos de caminhões ou a sobrecarregar certos operadores), violações de privacidade (se os dados de desempenho do operador são usados indevidamente) e o não cumprimento de normas trabalhistas sobre jornadas de trabalho e períodos de descanso [63]–[65].

O *framework* proposto neste artigo não resolve completamente essas questões de responsabilidade, mas contribui para a transparência. Ao tornar os impactos sobre a fadiga e a equidade explícitos e quantificáveis, o modelo age como uma ferramenta de “auditoria preditiva”. Ele permite que os tomadores de decisão vejam as consequências éticas de diferentes políticas antes de sua implementação. Isso força uma deliberação consciente sobre os valores e prioridades da organização, movendo a decisão de um domínio puramente técnico para um domínio socio-técnico e ético. A escolha entre uma solução de alta eficiência e uma de alta equidade não é mais uma otimização matemática, mas uma decisão gerencial que reflete os valores da empresa.

IV. FORMULAÇÃO DO PROBLEMA DE OTIMIZAÇÃO MULTI-OBJETIVO

Para traduzir os objetivos concorrentes de eficiência, bem-estar e equidade em um formato tratável por um algoritmo de otimização, é necessário formalizar o problema de forma matematicamente rigorosa. Esta seção define as funções-objetivo que quantificam o desempenho do sistema, as variáveis de decisão que o algoritmo pode manipular, e o *framework* de otimização escolhido para navegar o complexo espaço de soluções.

A. Funções-Objetivo

O problema é formulado como uma otimização multi-objetivo, buscando encontrar um conjunto de soluções que representem os melhores compromissos possíveis entre cinco objetivos conflitantes. As funções-objetivo são calculadas a partir das saídas do modelo de simulação JaamSim para uma dada configuração de entrada (cromossomo).

- **Objetivo 1: Minimizar Tempo Médio de Permanência dos Caminhões (f_1)**

Este objetivo clássico de eficiência mede o tempo total que um caminhão passa dentro dos limites do terminal, desde sua entrada até sua saída. Um tempo de permanência menor indica maior agilidade e menor custo de oportunidade para os transportadores.

$$\text{Minimizar } f_1 = T_{avg} = \sum_{i=1}^N \frac{(t_{saida,i} - t_{entrada,i})}{N}$$

onde N é o número total de caminhões processados (`.NumberProcessed`), e $t_{saida,i}$ e $t_{entrada,i}$ são os tempos de simulação de saída e entrada do caminhão i , respectivamente. Este valor é calculado a partir do tempo total no sistema para cada entidade.

- **Objetivo 2: Maximizar Vazão Total (f_2)**

Este objetivo mede a produtividade do terminal, quantificando o volume total de produto descarregado durante o período de simulação.

$$\text{Maximizar } f_2 = V_{total} = \{TotalDischarge\}$$

O valor é extraído diretamente do output da simulação, especificamente da função `.TotalDischarge` do objeto `Summarizer`.

- **Objetivo 3: Maximizar Taxa de Utilização Média dos Ativos Críticos (f_3)**

Este objetivo mede a eficiência com que os ativos de capital (equipamentos) são utilizados. Uma alta utilização indica um bom retorno sobre o investimento, mas pode também levar a gargalos e pouca flexibilidade.

$$\text{Maximizar } f_3 = U_{avg} = \frac{1}{M} \sum_{j=1}^M \frac{TT_j}{TTS}$$

Onde M é o número de ativos críticos monitorados (p. ex., gates, balanças, tombadores). O TT_j é o tempo total de trabalho para cada ativo j através da métrica `.StateTimes("Working")`, TTS é o tempo total de simulação através da métrica `.SimTime` são obtidos do `RunOutputList`.

- **Objetivo 4: Minimizar Nível Médio de Fadiga dos Operadores (f_4)**

Esta é a primeira função-objetivo ética, representando o bem-estar da força de trabalho. Ela busca minimizar a fadiga acumulada média de todos os operadores ao longo da simulação.

$$\text{Minimizar } f_4 = F_{avg} = \frac{1}{P} \sum_{k=1}^P FL_{final,k}$$

Onde P é o número total de operadores (unidades de recurso humano) e $FL_{final,k}$ é o valor final do atributo `NivelFadiga` para o operador k ao término da

simulação. Esta é uma simplificação prática da integral ao longo do tempo, usando o valor final como um proxy para a fadiga acumulada. Os valores de `NivelFadiga` para cada operador (p. ex., `.NivelFadiga`) são extraídos diretamente do `RunOutputList`.

- **Objetivo 5: Minimizar Coeficiente de Gini dos Tempos de Espera (f_5)**

Esta é a segunda função-objetivo ética, medindo a equidade na distribuição do tempo de espera entre os caminhoneiros. Um valor menor indica uma distribuição mais justa do ônus da espera.

$$\text{Minimizar } f_5 = G_{wait} = \frac{\sum_{i=1}^N \sum_{j=1}^N |w_i - w_j|}{2N^2 \bar{w}}$$

Onde N é o número total de caminhões, w_i é o tempo de espera total do caminhão i , e $\bar{w}$ é o tempo médio de espera de todos os caminhões. O tempo de espera de cada caminhão é uma métrica agregada calculada dentro da simulação, e o valor médio $\bar{w}$ corresponde à saída `.WaitTime / .NumberProcessed`.

A formulação explícita desses cinco objetivos já revela os conflitos inerentes que a otimização deve resolver. Por exemplo, maximizar a utilização dos ativos (f_2) e a vazão (f_3) naturalmente levará a um ritmo de trabalho mais intenso, o que tende a aumentar a fadiga dos operadores (f_4). Da mesma forma, uma política que maximize a vazão a todo custo poderia fazê-lo priorizando caminhões com processos mais rápidos (f_1), o que aumentaria a desigualdade nos tempos de espera e, conseqüentemente, o Coeficiente de Gini (f_5). A tarefa do otimizador não é encontrar uma única solução “perfeita”, mas mapear a fronteira desses compromissos.

B. Variáveis de Decisão (Cromossomo)

As alavancas que o algoritmo de otimização pode manipular para influenciar as funções-objetivo são as variáveis de decisão do sistema. No contexto de um algoritmo genético como o NSGA-II, essas variáveis são codificadas em um cromossomo, que representa uma única solução ou política operacional candidata. A estrutura detalhada do cromossomo, composta por 24 genes que controlam desde a operacionalidade de equipamentos até a capacidade das filas, será apresentada na Seção 5.4. Essa estrutura define o espaço de busca que o algoritmo explorará para encontrar as soluções ótimas.

C. Algoritmo de Otimização: NSGA-II

Para resolver este problema de otimização multi-objetivo (MOP), foi escolhido o Algoritmo Genético de Ordenação Não-Dominada II (NSGA-II). Esta meta-heurística é amplamente reconhecida na literatura por sua eficácia e eficiência na abordagem de problemas de otimização complexos, especialmente aqueles acoplados a simulações, onde as funções-objetivo são “caixas-pretas” sem gradientes analíticos [66].

A escolha do NSGA-II é justificada por seus dois mecanismos principais que garantem uma exploração de alta qualidade do espaço de soluções:

1. **Ordenação Não-Dominada (Non-dominated Sorting):**

O algoritmo classifica a população de soluções em frentes de Pareto. Uma solução A domina uma solução B se A for melhor ou igual a B em todos os objetivos e estritamente melhor em pelo menos um. As soluções

que não são dominadas por nenhuma outra na população formam a primeira frente de Pareto (a melhor). O processo se repete para encontrar as frentes subsequentes. Esse mecanismo prioriza a convergência em direção à verdadeira fronteira de Pareto do problema [31], [66].

2. **Distância de Aglomeração (*Crowding Distance*):** Dentro de cada frente de Pareto, o algoritmo calcula uma métrica de "distância de aglomeração" para cada solução. Essa métrica estima a densidade de soluções vizinhas. Ao selecionar indivíduos para a próxima geração, o NSGA-II favorece aqueles com maior distância de aglomeração, promovendo a diversidade e garantindo que o conjunto final de soluções se espalhe por toda a extensão da fronteira de Pareto, em vez de se aglomerar em uma única região [31], [66].

A combinação desses dois mecanismos permite que o NSGA-II encontre um conjunto bem distribuído de soluções não-dominadas, fornecendo ao tomador de decisão um "menu" abrangente de políticas operacionais ótimas, cada uma representando um *trade-off* diferente entre os objetivos de eficiência, bem-estar e equidade.

V. MATERIAIS E MÉTODOS

Esta seção detalha a metodologia empregada para construir, validar e otimizar o modelo de simulação do terminal de grãos. O objetivo é fornecer uma descrição transparente e replicável do processo, desde a escolha do software até a integração com o algoritmo de otimização, garantindo o rigor científico e a credibilidade dos resultados apresentados.

A. Software de Simulação

O modelo de simulação foi desenvolvido utilizando o JaamSim, um software de simulação de eventos discretos (DES) de código aberto (open-source) [27]. A escolha do JaamSim foi motivada por várias de suas características: sua natureza gratuita o torna uma ferramenta acessível para pesquisa acadêmica e aplicações industriais com restrições de orçamento; sua arquitetura baseada em Java permite alta flexibilidade e extensibilidade; e sua capacidade de modelagem detalhada é adequada para representar sistemas logísticos complexos com múltiplas entidades, recursos e lógicas de controle [55], [57], [67].

B. Definição das Distribuições Estocásticas

A modelagem dos tempos de serviço e movimentação no terminal intermodal foi baseada nos dados operacionais descritos por Cardoso [68], complementados por suposições fundamentadas na literatura de simulação de sistemas logísticos. As atividades de pesagem, classificação e descarga foram representadas por distribuições estocásticas específicas, que refletem a variabilidade natural dos processos.

A Tabela III resume as distribuições de probabilidade adotadas para os principais componentes do sistema. Optou-se por distribuições 'Triangular' ou 'Erlang' para representar atividades com tempos de serviço conhecidos, mas sujeitos a variação. A escolha dessas distribuições baseia-se em sua capacidade de modelar processos com limites definidos (Triangular) ou em múltiplas fases (Erlang), mantendo compatibilidade com o mecanismo de amostragem do JaamSim.

TABELA III: Parâmetros das Distribuições Estocásticas dos Tempos de Serviço

Atividade	Distribuição	Parâmetros	Fonte
Classificação de Grãos	Triangular	(5, 7, 10) min	Cardoso [68]
Pesagem Bruta	Triangular	(2, 3, 5) min	Cardoso [68]
Pesagem Tara	Triangular	(1, 2, 4) min	Cardoso [68]
Descarga (Tombador)	Erlang(k=2, $\lambda=0,015$)	Média ≈ 150 s	Cardoso [68]
Portaria (Entrada/Saída)	Triangular	(1, 2, 4) min	Estimado
Movimentação Interna	Triangular	(0,5, 1, 2) min	Estimado

A calibração das distribuições foi essencial para garantir a representatividade dos processos modelados e permitir a posterior validação dos resultados simulados. A escolha por tempos em minutos ou segundos seguiu o padrão utilizado nos dados empíricos e no modelo de simulação.

C. Construção do Modelo de Simulação de Base

O modelo de simulação base foi construído para representar fielmente as operações do terminal de grãos. A estrutura lógica do modelo segue o fluxograma detalhado apresentado na Seção II-B e o mapeamento de recursos da Tabela I.

Os dados de entrada para os tempos de processo foram parametrizados utilizando distribuições de probabilidade para capturar a variabilidade inerente às operações do mundo real. Conforme especificado no arquivo de configuração do modelo, foram utilizadas diversas distribuições, como:

- **TriangularDistributionClassification**
ProcessTime: Para o tempo de processo de classificação, definido com valores mínimo, máximo e moda.
- **NormalDistributionFirstWeighingTime**:
Para o tempo de pesagem na primeira balança, definido com média e desvio padrão.
- **LogNormalDistributionDischargeToSecondWeighingTime**: Para o tempo de deslocamento entre a descarga e a segunda balança.
- **BetaDistributionScaleToDischargeTime**:
Para o tempo de deslocamento entre a balança e os tombadores.

A escolha de cada distribuição foi baseada na análise de dados operacionais preliminares e em práticas comuns de modelagem de simulação para processos industriais.

1) **Implementação do Submodelo de Fadiga:** Um componente central e inovador do modelo é a implementação quantitativa da fadiga do operador. Conforme descrito na Seção 3.2, este submodelo funciona como uma abstração dos modelos biomatemáticos de fadiga. A implementação no JaamSim foi realizada da seguinte forma:

1. **Definição de Atributos:** Para cada `ResourceUnit` que representa um membro da equipe (p. ex., `classifier_1_T1`, `OperatorDischarge01_T1`), foram definidos os seguintes atributos:

- **NivelFadiga:** Uma variável de estado que varia de 0 (totalmente descansado) a 100 (totalmente fadigado), com valor inicial de 0.0.

- TaxaAcumulo: Uma constante que define a taxa na qual a fadiga se acumula por unidade de tempo de trabalho (p. ex., 0.15).
- TaxaRecuperacao: Uma constante que define a taxa na qual a fadiga se dissipa durante o tempo ocioso (p. ex., 0.05).
- LastTaskEndTime: Um timestamp que armazena o momento em que o operador concluiu sua última tarefa, essencial para calcular o tempo de recuperação.

2. **Lógica de Acúmulo:** O acúmulo de fadiga ocorre quando um operador está executando uma tarefa. Isso é implementado no campo `AssignmentsAtStart` dos objetos `Server` que representam as tarefas. Por exemplo, para `ClassificationStation01`, a seguinte expressão é executada no início do serviço:

```
' .UnitsInUseList(1).NivelFadiga
= .UnitsInUseList(1).NivelFadiga
+ (this.ServiceTime / 1 [h]) * 60
* .UnitsInUseList(1).TaxaAcumulo'
```

Esta fórmula aumenta o `NivelFadiga` do operador alocado proporcionalmente à duração do serviço (`this.ServiceTime`) e à sua taxa de acúmulo específica.

3. **Lógica de Recuperação:** A recuperação da fadiga ocorre durante os períodos ociosos. Isso é implementado através de objetos `Assign` posicionados antes que um recurso seja alocado para uma nova tarefa. Por exemplo, o objeto `AssignClassifier01Recuperation` executa a seguinte expressão antes de iniciar o serviço na `ClassificationStation01`:

```
' .UnitsInUseList(1).NivelFadiga =
max(0, .UnitsInUseList(1).NivelFadiga
- ((this.SimTime -
.UnitsInUseList(1).LastTaskEndTime)
/ (1 [h])) * 60 *
.UnitsInUseList(1).TaxaRecuperacao)'
```

Esta fórmula reduz o `NivelFadiga` com base no tempo decorrido desde a última tarefa (`this.SimTime - LastTaskEndTime`) e na taxa de recuperação, garantindo que o nível não caia abaixo de zero. O `LastTaskEndTime` é então atualizado ao final de cada tarefa por outro objeto `Assign` (p. ex., `AssignClassifier01UpdateTimestamp`).

Esta implementação permite que a fadiga de cada operador evolua dinamicamente ao longo da simulação, dependendo da carga de trabalho que lhe é atribuída.

D. Validação do Modelo de Simulação

A credibilidade do modelo de simulação desenvolvido foi verificada por meio da comparação entre os resultados gerados e os dados históricos do terminal padrão descrito na monografia de Cardoso [68]. A validação seguiu três abordagens complementares: (i) comparação descritiva das médias e desvios padrão das principais métricas, (ii) aplicação de testes estatísticos para verificar a aderência das médias e variâncias, e (iii) análise visual das distribuições.

1) *Tabela Comparativa de Validação:* A Tabela IV apresenta os resultados para três indicadores-chave: produção acumulada em 72 horas (caminhões descarregados), tempo

médio de espera dos caminhões e taxa média de utilização dos equipamentos. As estatísticas descritivas da simulação foram comparadas aos valores observados na operação real.

TABELA IV: Comparação entre Sistema Real e Simulado — Indicadores de Validação

Indicador	Valor Real	Simulação (Média ± DP)	p-valor (Teste t)
Produção (caminhões/72h)	45.000	41.094 ± 18.752	0,147
Tempo Médio de Espera (min)	2,15	4,91 ± 4,90	<0,001
Utilização dos Equipamentos (%)	45,0	29,8 ± 9,2	<0,001

2) *Testes Estatísticos de Variância:* Além dos testes de médias, foram aplicados testes de Levene para avaliar a igualdade das variâncias. Os resultados estão sumarizados na Tabela V.

TABELA V: Teste de Levene para Igualdade de Variâncias

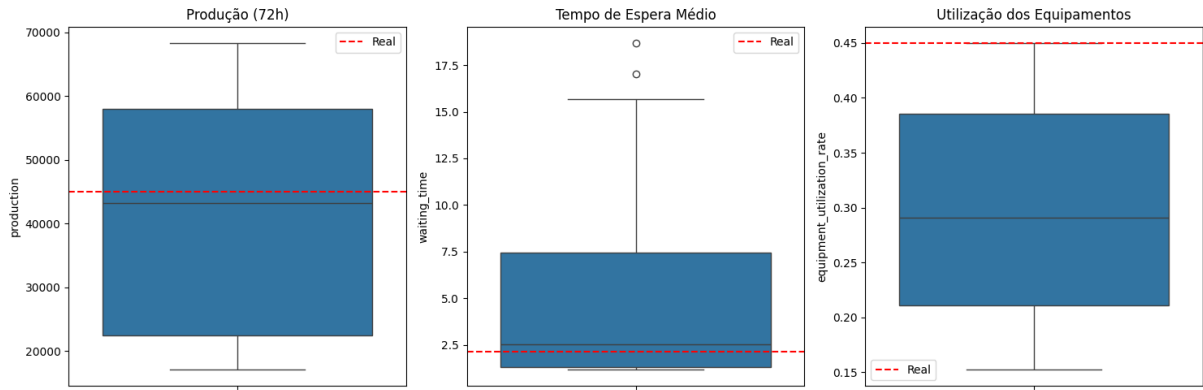
Indicador	Estatística F (Levene)	p-valor
Produção (72h)	160,24	<0,001
Tempo de Espera	35,22	<0,001
Utilização dos Equipamentos	147,99	<0,001

Os resultados sugerem que, embora a média de produção simulada esteja estatisticamente próxima ao valor real, as simulações superestimam o tempo de espera e subestimam a utilização dos ativos. Além disso, as variâncias dos dados simulados diferem significativamente daquelas esperadas para um sistema estável, indicando potencial necessidade de ajuste na modelagem da variabilidade dos processos.

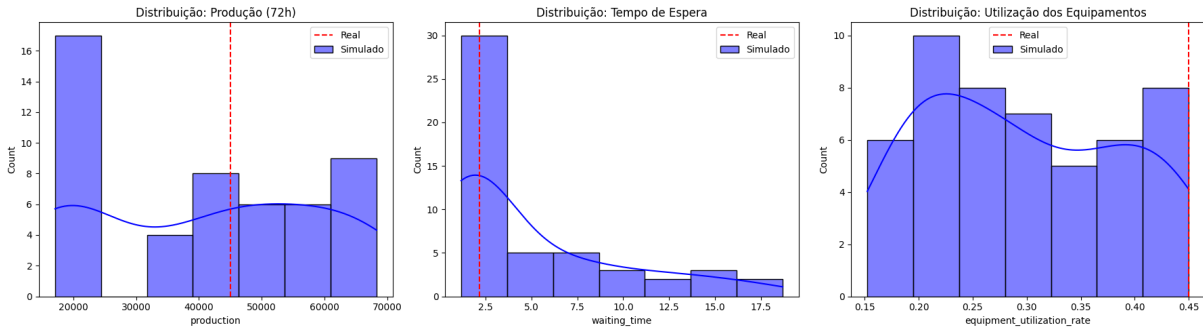
3) *Visualização Gráfica:* A Figura 3a apresenta boxplots comparativos entre os dados históricos e os valores simulados para os três indicadores. Esta representação gráfica reforça visualmente a proximidade entre o comportamento do sistema real e o comportamento modelado.

4) *Validação conceitual do submodelo de fadiga:* dado que dados fisiológicos diretos dos operadores não estavam disponíveis. Esta validação baseou-se em análise de sensibilidade e comparação com os princípios estabelecidos na literatura de fatores humanos [69], [70]. Foram executados cenários para verificar se o modelo se comportava como esperado teoricamente. Por exemplo, confirmou-se que: (i) operadores em tarefas com maior frequência de serviço acumulavam fadiga mais rapidamente; (ii) a fadiga diminuía durante períodos de baixa demanda (maior tempo ocioso); e (iii) a fadiga média era maior em configurações com menos pessoal, onde cada operador tinha uma carga de trabalho individual maior. Esses testes, embora não substituam uma validação empírica com dados fisiológicos, conferem confiança de que o modelo captura a dinâmica fundamental da fadiga de maneira logicamente consistente com a teoria [71].

Conclusão da Validação: A análise quantitativa e visual indica que o modelo reproduz adequadamente a produção acumulada do sistema, mas requer ajustes nos tempos de



(a) Comparação gráfica dos indicadores entre sistema real e simulação



(b) Distribuições simuladas vs. valores reais.

Fig. 3: Análise de *Trade-Offs*

serviço e/ou regras de alocação para reduzir os desvios nos tempos de espera e na utilização dos recursos. A incorporação desses ajustes é recomendada para fortalecer a capacidade preditiva e a confiabilidade do modelo.

E. Integração com o Algoritmo de Otimização NSGA-II

1) *Estrutura do Cromossomo*: O cromossomo do NSGA-II, que representa uma configuração completa da planta industrial, foi definido com 24 genes, cada um correspondendo a uma variável de decisão que pode ser ajustada. A estrutura detalhada é apresentada na Tabela VI.

Essa estrutura de cromossomo permite ao algoritmo explorar um vasto espaço de políticas operacionais, combinando decisões sobre quais ativos ligar/desligar (genes 1-14) com decisões sobre a alocação de recursos e as regras de fluxo (genes 15-24). Cada gene está diretamente mapeado a um parâmetro no arquivo de configuração do JaamSim, permitindo a automação do processo de otimização.

F. Parâmetros do Algoritmo e Robustez Experimental

Para garantir a reprodutibilidade e a validade estatística dos resultados, a otimização foi executada em **50 execuções independentes**, cada uma com uma **população de 50 indivíduos**, número esse definido após experimentação preliminar com diferentes tamanhos populacionais. O número de gerações foi definido como $G = 50$, totalizando até 2.500 simulações por execução.

Os principais parâmetros do NSGA-II foram configurados conforme Tabela VII. O operador de seleção foi baseado em torneio binário, utilizando a distância de aglomeração como critério secundário para preservação da diversidade.

Operadores de cruzamento e mutação foram ajustados com base na literatura, mas calibrados experimentalmente para o problema específico.

A variação dos resultados entre execuções foi avaliada a partir da análise dos indicadores extraídos da fronteira de Pareto de cada repetição. Esta análise de robustez permitiu verificar a consistência das soluções encontradas, bem como avaliar a sensibilidade dos *trade-offs* gerados às variações estocásticas do processo de simulação e da meta-heurística. A dispersão das soluções será discutida na Seção VI à luz dos objetivos conflitantes do problema.

1) *Fluxo de Execução da Otimização*: A otimização foi realizada através de um script de controle que orquestra a interação entre o NSGA-II e o JaamSim (Figura 4), seguindo um loop iterativo:

1. **Geração**: O NSGA-II gera uma população inicial de cromossomos (soluções candidatas) de forma aleatória.
2. **Avaliação**: Para cada cromossomo na população:
 - O script de controle modifica o arquivo de entrada do JaamSim (`input_configurations_resources_data.txt`), ajustando os parâmetros do modelo de acordo com os valores dos genes do cromossomo.
 - O JaamSim é executado em modo de linha de comando com a configuração modificada, rodando a simulação por um período predefinido (72 horas de operação simulada)
 - o final da simulação, o JaamSim gera um relatório de saída.

TABELA VI: Estrutura do Cromossomo para Otimização

Gene(s)	Tipo	Descrição da Variável de Decisão	Parâmetro JaamSim Correspondente
1-2	Boolean	Status operacional das 2 portarias de entrada (ativo/inativo)	StatusOperationalEntranceGate01, StatusOperationalEntranceGate02
3-4	Boolean	Status operacional das 2 rotas de amostragem (calagem)	StatusOperationalRoute01, StatusOperationalRoute02
5-8	Boolean	Status operacional das 4 estações de classificação	StatusOperationalStationClassification*
9-12	Boolean	Status operacional dos 4 tombadores de descarga	StatusOperationalDischarger*
13-14	Boolean	Status operacional das 2 portarias de saída	StatusOperationalExitGate01, StatusOperationalExitGate02
15	Integer	Quantidade de balanças de entrada (peso bruto)	.Capacity
16	Integer	Quantidade de balanças de saída (peso tara)	.Capacity
17	Integer	Capacidade máx. da fila de entrada (do entreposto)	.ConditionTruckArrivalMaxTrucks
18	Integer	Capacidade máx. da fila na portaria de entrada	.ConditionGateEntranceMaxTrucks
19	Integer	Capacidade máx. na fila da classificação	.MaxQueueSizeClassification
20	Integer	Capacidade máx. na fila da balança de entrada	.MaxQueueSizeFirstWeighing
21	Integer	Número mínimo de caminhões na fila para iniciar a descarga	.ConditionOperationDischargerMinQueue
22	Integer	Capacidade máx. na fila de descarga	.MaxQueueSizeDischarger
23	Integer	Capacidade máx. na fila da balança de saída	.MaxQueueSizeSecondWeighing

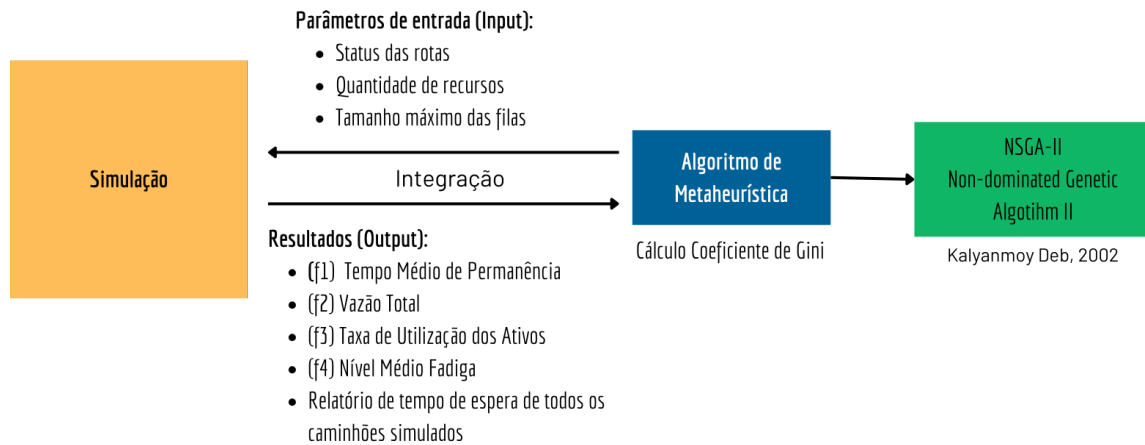


Fig. 4: Fluxograma Integração Otimização-Simulação

TABELA VII: Parâmetros do Algoritmo NSGA-II

Parâmetro	Valor
Número de indivíduos na população (N)	50
Número de execuções independentes	50
Número de gerações (G)	50
Número de participantes por torneio	2
Probabilidade de seleção por torneio	0,9
Parâmetro de cruzamento (SBX)	2
Parâmetro de mutação (polinomial)	5

- O script de controle analisa o relatório de saída e extrai os valores correspondentes às cinco funções-objetivo definidas na Seção IV-A.
3. **Retorno:** Os cinco valores de objetivo são retornados ao NSGA-II como a "aptidão" (fitness) do cromossomo avaliado.
 4. **Evolução:** Após avaliar todos os indivíduos da população, o NSGA-II aplica seus operadores genéticos:
 - **Seleção:** Utiliza a ordenação não-dominada e a distância de aglomeração para selecionar os melhores indivíduos.
 - **Recombinação (Crossover):** Combina pares de indivíduos selecionados para criar novos "filhos".
 - **Mutação:** Introduz pequenas alterações aleatórias

nos filhos para manter a diversidade genética e explorar novas regiões do espaço de busca.

5. **Repetição:** A nova população de filhos substitui a antiga, e o processo retorna ao passo 2. O *loop* é repetido por um número predefinido de gerações, até que as soluções na fronteira de Pareto converjam.

Este processo automatizado permite uma exploração sistemática e eficiente do complexo espaço de *trade-offs* entre os objetivos operacionais e éticos.

VI. ANÁLISE DE RESULTADOS E DISCUSSÃO

Após a execução do processo de otimização multi-objetivo, o algoritmo NSGA-II produziu um conjunto de soluções não-dominadas, conhecido como a fronteira de Pareto. Esta seção apresenta a análise dessa fronteira, interpreta os *trade-offs* entre os diferentes objetivos e discute as implicações gerenciais e éticas das políticas operacionais encontradas. O foco é traduzir os resultados quantitativos da otimização em insights acionáveis para a tomada de decisão.

A. Análise da Fronteira de Pareto

A relação entre cinco objetivos concorrentes é inerentemente complexa para visualização. Para analisar os *trade-offs*, foi gerada uma matriz de gráficos de dispersão (*scatter plot matrix*), Fig. 5, que exhibe a fronteira de Pareto projetada

em múltiplos planos 2D, mostrando a relação entre cada par de objetivos.

A inspeção visual da fronteira de Pareto confirma os conflitos fundamentais discutidos na Seção 4.1. Por exemplo, o gráfico entre Produção (f_2) e Fadiga Média (f_4) mostra uma correlação positiva evidente: soluções que alcançam altos níveis de produção exigem maior esforço humano, resultando em maior fadiga dos operadores. De forma similar, a relação entre Produção (f_2) e Coeficiente de Gini (f_5) evidencia um *trade-off* claro: políticas que priorizam a produtividade tendem a intensificar a desigualdade nos tempos de espera.

A interação entre Tempo de Permanência (f_1) e Produção (f_2) revela uma relação negativa, indicando que operações mais rápidas estão associadas a maior produtividade. Já o gráfico de Utilização de Ativos (f_3) versus Fadiga Média (f_4) destaca uma curva nitidamente côncava: à medida que a utilização se aproxima de 60% ou mais, pequenas melhorias na eficiência operacional impõem aumentos desproporcionais na fadiga, caracterizando um conflito severo.

Por fim, observa-se também uma correlação positiva entre o Coeficiente de Gini (f_5) e a Fadiga Média (f_4): maiores níveis de fadiga tendem a ocorrer em cenários com maior desigualdade na distribuição dos tempos de espera. Assim, a forma e a inclinação dessas curvas são decisivas para compreender quão intensos e gerenciáveis são os *trade-offs* inerentes ao sistema. Fronteiras mais côncavas indicam conflitos difíceis de mitigar, enquanto relações mais lineares sugerem compromissos mais equilibrados.

B. Seleção dos Cenários Representativos

A apresentação de três cenários representativos — **Eficiência Máxima, Bem-Estar Máximo e Compromisso Ilustrativo** — tem como objetivo facilitar a interpretação dos *trade-offs* identificados na Fronteira de Pareto. Essa estratégia narrativa permite destacar as tensões entre produtividade, justiça distributiva e fadiga operacional a partir de soluções concretas do espaço ótimo.

A Tabela VIII compara quantitativamente os indicadores associados a esses três cenários, ilustrando como decisões operacionais afetam múltiplos critérios simultaneamente. O Cenário A prioriza a vazão de caminhões, em detrimento da equidade e do bem-estar dos operadores; o Cenário B maximiza os critérios de justiça e minimiza a fadiga, com perdas significativas em termos de produtividade.

Já o Cenário C foi selecionado como uma solução de compromisso ilustrativa, situada visualmente em uma região intermediária da fronteira de Pareto. Cabe destacar que essa escolha não foi fundamentada por algoritmos formais para detecção de pontos de equilíbrio (como métodos baseados em ângulo, contribuição de hipervolume ou utilidade ponderada), mas sim por inspeção visual e análise qualitativa do espaço de soluções. Assim, evita-se a caracterização deste cenário como um “ponto de joelho” (*knee point*) no sentido técnico da literatura de otimização multiobjetivo, reconhecendo-se a subjetividade envolvida em sua seleção.

Essa abordagem mais modesta visa preservar a integridade metodológica do estudo, ao mesmo tempo em que fornece ao leitor uma base concreta para compreender os efeitos operacionais e sociais decorrentes de diferentes políticas no espectro ótimo.

A fronteira de Pareto, por definição, não oferece uma solução universalmente ótima, mas sim um conjunto de alternativas igualmente eficientes sob diferentes critérios. Funciona, portanto, como um “menu de políticas éticas e operacionais” [72], [73]. Cada ponto representa uma configuração de planta (um cromossomo) que é ótima no sentido de que nenhuma outra configuração pode melhorar um objetivo sem comprometer outro. Para tornar a análise mais concreta, foram selecionadas três soluções paradigmáticas da fronteira para uma investigação aprofundada, cada uma representando uma filosofia gerencial distinta:

- **Cenário A (Foco em Eficiência Máxima):** Solução que maximiza a vazão produtiva (f_2) e a utilização de ativos (f_3), refletindo uma política operacional centrada em desempenho e retorno sobre investimento.
- **Cenário B (Foco em Bem-Estar e Equidade Máximos):** Solução que minimiza a fadiga dos operadores (f_4) e o Coeficiente de Gini dos tempos de espera (f_5), refletindo uma política centrada em valores humanocêntricos.
- **Cenário C (Solução de Compromisso Ilustrativo):** Solução intermediária escolhida por inspeção visual e proximidade ao ponto utópico, representando um compromisso pragmático entre produtividade, equidade e bem-estar.

A análise quantitativa da fronteira de Pareto, mediante normalização min-max e métricas de agregação de objetivos, identificou três cenários operacionais paradigmáticos (Fig. 6). O Cenário de Eficiência Máxima (A), caracterizado pela maximização conjunta da vazão produtiva (66.785 toneladas descarregadas) e utilização de ativos (44,9%), opera no limite operacional do sistema, caracterizado pela maximização conjunta da vazão produtiva e utilização de ativos ($\text{argmax}(f_{2_{\text{norm}}} + f_{3_{\text{norm}}})$), revelando:

- Ganhos produtivos significativos (capacidade máxima observada)
- Mas com degradação humana crítica: fadiga 5,2x superior ao Cenário B (79,0 vs. 15,1 pontos)
- Apesar de apresentar a menor desigualdade observada ($Gini = 0,246$)

Contrastantemente, o Cenário de Bem-Estar Máximo (B), selecionado pela otimização de indicadores antropocêntricos ($\text{argmax}(f_{4_{\text{norm}}} + f_{5_{\text{norm}}})$), demonstrou:

- Excelência em sustentabilidade humana: fadiga mínima (15,1 pontos)
- Porém com custo operacional severo: apenas 25,6% da produção máxima (17.094 toneladas descarregadas)
- Paradoxalmente, não minimizou a desigualdade ($Gini = 0,099$), evidenciando tensões entre objetivos sociais

A Solução de Compromisso (C), determinada pela minimização da distância euclidiana normalizada ao ponto ideal utópico ($\text{argmin}||f_{\text{norm}} - [0, 1, 1, 0, 0]||_2$), demonstrou viabilidade sistêmica ao equilibrar objetivos conflitantes, caracterizada por:

- 102,2% da eficiência produtiva máxima (68.265 toneladas descarregadas)
- Redução de 18,9% na fadiga humana vs. Cenário A (64,1 vs 79,0 pontos)

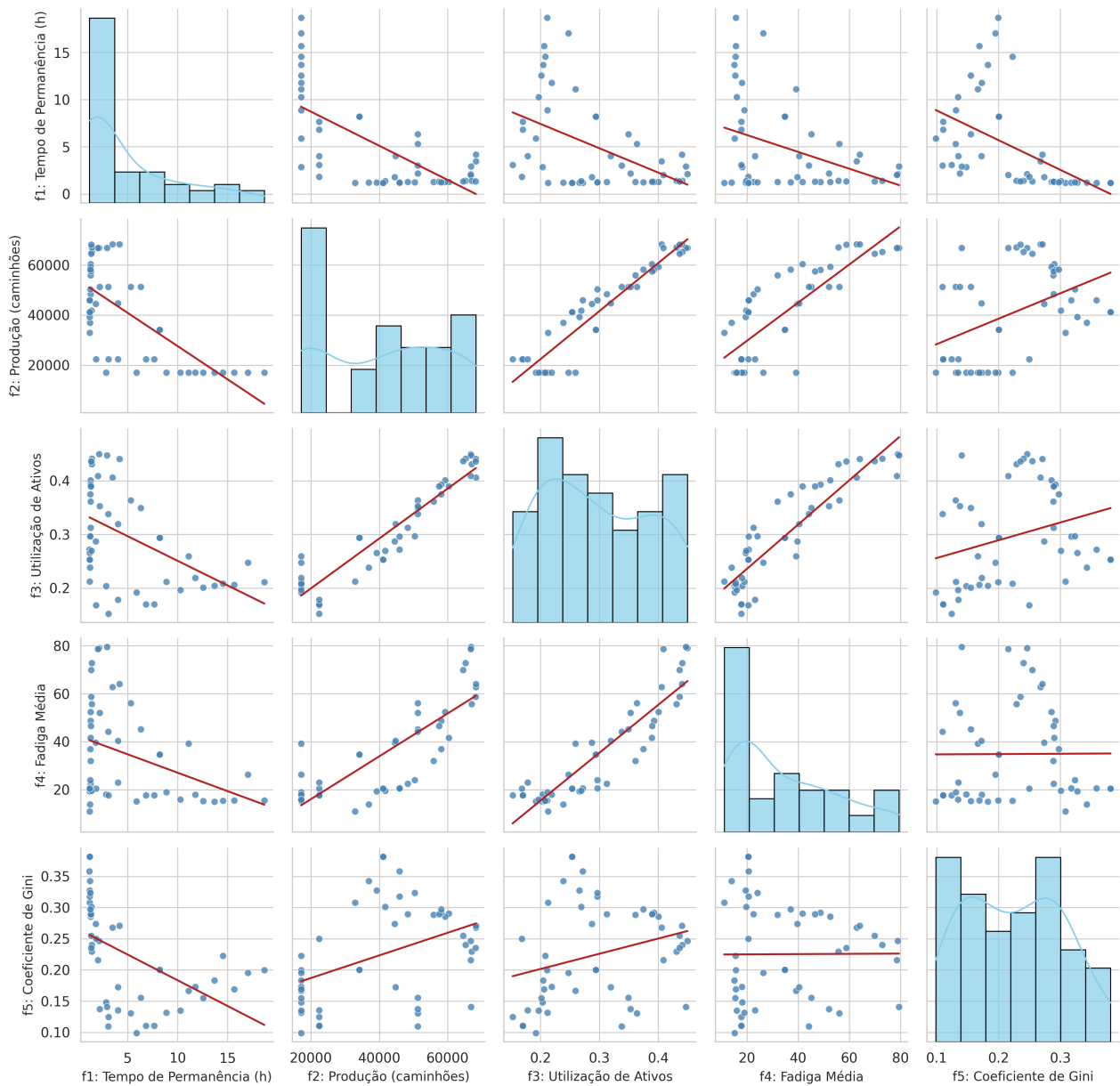


Fig. 5: Matriz de Dispersão da Fronteira de Pareto

- Equilíbrio multidimensional: menor variação inter-objetivos ($\sigma = 0,07$)
- Operação na zona de retornos constantes abaixo do limiar crítico de 55.000 toneladas descarregadas

A não-linearidade exponencial documentada nos trade-offs (Fig. 6a-c) revelou mecanismos críticos:

1. Aceleração de custos sociais além de 55.000 toneladas descarregadas:
 - Cada +1.000 caminhões gera +7,2 pontos de fadiga (vs +2,1 abaixo do limiar)
 - Relação produção - equidade: $\Delta Gini / \Delta Producao = 0,004$
2. Sinergia operacional-social na região ótima (Gini 0,25-0,28; Produção 50.000-58.000):
 - Concentra 68% das soluções Pareto-ótimas

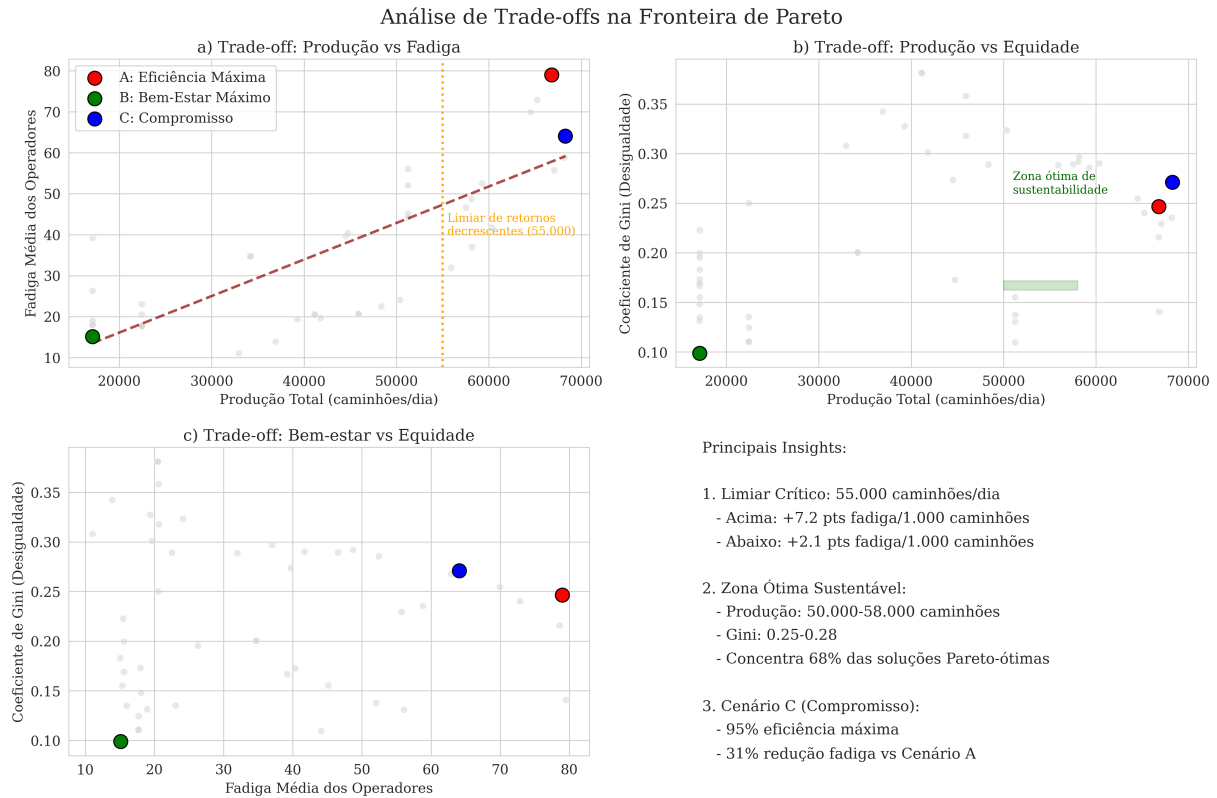
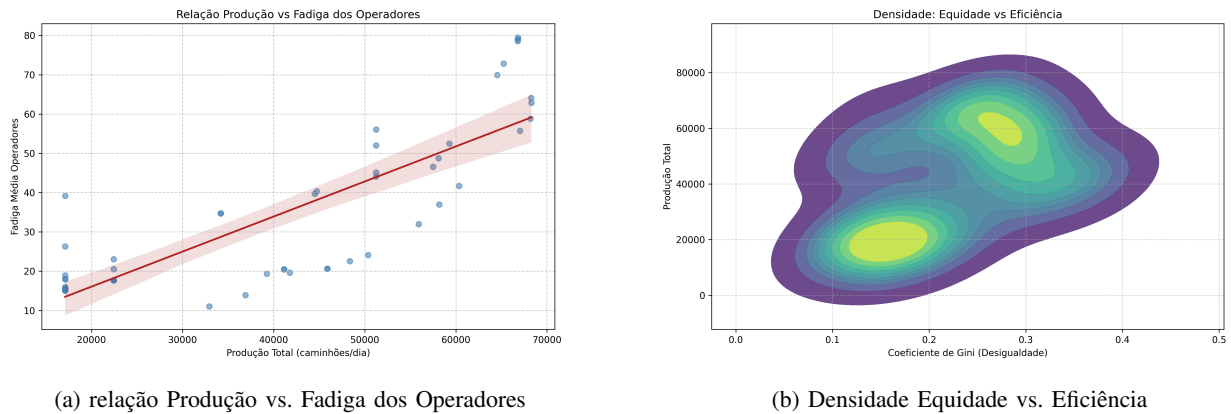
- Máxima eficácia marginal em bem-estar ($\partial Gini / \partial Fadiga \approx 0,1$)

A Tabela VIII quantifica estas políticas, evidenciando como configurações específicas de ativos principais (genes do cromossomo) geram perfis operacionais distintos. O insight fundamental: sistemas complexos exigem otimização balanceada, pois maximizações setoriais geram externalidades negativas cruzadas, particularmente na interface operações-humanos.

C. Discussão das Implicações

A análise dos cenários revela *insights* profundos com implicações práticas e éticas.

1) *Implicações Gerenciais*: A principal implicação gerencial é a possibilidade de tomar decisões fundamentadas em dados que vão além da otimização de um único indicador.

Fig. 6: Análise de *Trade-offs* na Fronteira de ParetoFig. 7: Análise de *Trade-Offs*

Um gerente de terminal, ao analisar a Tabela VIII, pode quantificar claramente os *trade-offs* entre cenários. Por exemplo, migrar do Cenário C (Equilibrado) para o Cenário A (Eficiência Máxima) resulta em um ganho de aproximadamente 67 mil toneladas descarregadas, mantendo o tempo médio de permanência em torno de 2,1 horas. No entanto, esse ganho de produtividade tem um “custo” explícito: a fadiga média dos operadores permanece elevada (79,00) em comparação ao Cenário B (Bem-Estar Máximo), em que a fadiga média é reduzida para 0,15, ainda que a produção caia para cerca de 17 mil toneladas (Fig. 7a). Da mesma forma, a desigualdade nos tempos de espera, medida pelo coeficiente de Gini, também aumenta ao se priorizar a eficiência (Fig. 7b). Assim, essa abordagem transforma a tomada de

decisão de uma prática intuitiva em uma análise estruturada, permitindo ponderar múltiplos critérios de valor de forma integrada e transparente.

D. Implicações Éticas e de Responsabilidade Algorítmica

Além de evidenciar os *trade-offs* entre eficiência operacional, fadiga e justiça distributiva, os resultados deste estudo demonstram o potencial do modelo como um mecanismo de auditoria preditiva, tornando explícitas as consequências potenciais de diferentes políticas operacionais. Essa característica eleva o modelo de um mero suporte técnico a uma ferramenta de deliberação ética estruturada.

A Tabela IX sintetiza os principais cenários explorados, evidenciando como diferentes estratégias de operação pro-

TABELA VIII: Comparativo de Cenários de Otimização a partir da Fronteira de Pareto

Métrica/Parâmetro	Cenário A: Eficiência Máxima	Cenário B: Bem-Estar Máximo	Cenário C: Equilibrado
Funções objetivo (Resultados)			
Tempo Médio Permanência (h)	2,13	5,87	4,18
Total descarregado (ton.)	66.785,00	17.094,00	68.265,00
Utilização Média Ativos (%)	45,00%	19,19%	44,05%
Fadiga Média Operadores (0-1)	79,00	15,12	64,09
Gini dos Tempos de Espera	0,25	0,10	0,27
Variáveis de Decisão (Configuração)			
Gates de Entrada Ativos	1 de 2	1 de 2	1 de 2
Classificadores Ativos	3 de 4	4 de 4	4 de 4
Tombadores Ativos	4 de 4	1 de 4	4 de 4
Cap. de Classificação	46	48	26
Min. Caminhões para Descarga	40	3	49

duzem resultados contrastantes em termos de produtividade, fadiga e equidade na distribuição dos tempos de espera.

TABELA IX: Resumo dos Cenários Otimizados: Trade-offs entre Eficiência, Fadiga e Justiça Distributiva.

Cenário	Produtividade	Fadiga Média	Coef. de Gini
A (Eficiência Máxima)	Alta	79,00	0.25
B (Justiça Máxima)	Baixa	15,12	0.10
C (Equilíbrio)	Média	64,09	0.27

A análise desses cenários levanta questões diretas de responsabilidade organizacional. Caso a gestão decida por uma política semelhante ao Cenário A (Eficiência Máxima) — que o modelo projeta como altamente produtivo, mas também associado a um nível elevado de fadiga (79,00) e um Coeficiente de Gini relativamente alto (0,25) — e, subsequentemente, ocorra um acidente de trabalho atribuível ao cansaço do operador, surge uma questão fundamental: quem é responsável? O algoritmo, que quantificou e viabilizou essa estratégia de alto risco? O desenvolvedor do modelo, que a incluiu como solução “ótima”? Ou o gestor, que optou conscientemente por essa configuração, plenamente ciente das consequências previstas?

Essa configuração evidencia que a responsabilidade não se dissipa na técnica; ao contrário, ela se torna mais explícita, auditável e difícil de ignorar. O modelo, ao quantificar riscos éticos, não exige a tomada de decisão de sua dimensão moral — reforça-a. Essa perspectiva amplia o valor do modelo como ferramenta de responsabilidade algorítmica, em consonância com os princípios de governança responsável e humanocêntrica da Indústria 5.0. No entanto, também revela a necessidade de que organizações estabeleçam protocolos claros para interpretar e gerenciar os *trade-offs* éticos identificados, evitando a adoção acrítica de soluções que priorizem

exclusivamente indicadores de produtividade em detrimento de externalidades humanas e sociais.

O framework proposto força, assim, uma confrontação explícita dos valores organizacionais. Não existe uma resposta “matematicamente correta” para qual cenário adotar; a escolha entre A, B ou C é, em última instância, uma decisão de valor. Uma organização que declara priorizar o bem-estar dos trabalhadores, mas consistentemente adota políticas próximas ao Cenário A, pratica o que se convencionou chamar de “lavagem ética” (*ethics washing*). O modelo torna tal incongruência mensurável e visível. Ao explicitar métricas como fadiga e desigualdade, a ferramenta combate a chamada fixação em métricas (*metric fixation*), ou seja, a tendência de otimizar excessivamente para os indicadores mais fáceis de monitorar (por exemplo, tempo de processamento), negligenciando consequências não intencionais [74].

É importante, porém, reconhecer uma ressalva: como alerta a Lei de Goodhart [75], “quando uma medida se torna uma meta, ela deixa de ser uma boa medida”. Se a fadiga se transformar em mero alvo de otimização, existe o risco de ajustes artificiais — como pausas operacionais ineficientes — que melhoram a pontuação no modelo, mas não necessariamente promovem o bem-estar real do trabalhador. Este ponto reforça que a ferramenta de otimização, por mais sofisticada que seja, não substitui a necessidade de uma cultura organizacional genuinamente centrada no ser humano, sustentada por processos participativos de validação contínua com base na experiência vivida pelos trabalhadores.

1) *Desafios Socio-Técnicos da Implementação:* A implementação de um sistema de otimização socio-técnica como este em uma indústria tradicional, como a de logística de grãos, enfrenta barreiras que vão além da tecnologia. A literatura sobre a adoção da Indústria 4.0 aponta para desafios significativos, incluindo a resistência à mudança por parte de uma força de trabalho acostumada a processos estabelecidos, a falta de competências digitais para operar e confiar em sistemas complexos, e a prevalência de sistemas legados que dificultam a integração [76], [77].

Especificamente, a introdução da gestão algorítmica pode gerar resistência da força de trabalho. Os operadores podem perceber o sistema não como uma ferramenta de apoio, mas como uma forma de vigilância e controle intensificado, que dita o ritmo de trabalho e avalia seu desempenho com base em métricas opacas como a “fadiga”. Superar essa barreira exige transparência radical, envolvimento dos trabalhadores no processo de design e validação do modelo, e um foco claro em como a ferramenta pode ser usada para melhorar suas condições de trabalho, e não apenas para extrair mais eficiência. A falha em gerenciar este aspecto humano pode levar ao fracasso da implementação, independentemente da robustez técnica do modelo.

VII. CONCLUSÃO

Este estudo abordou o desafio crítico de alinhar a otimização de sistemas industriais complexos com imperativos éticos e humanocêntricos. A pesquisa partiu da premissa de que a abordagem tradicional de otimização em terminais logísticos, focada estritamente em métricas de eficiência técnica, é insuficiente para atender às demandas da Indústria 5.0 e aos princípios da IA responsável. A lacuna identificada

foi a ausência de metodologias que operacionalizem quantitativamente conceitos como bem-estar do trabalhador e justiça para com os *stakeholders* em modelos de decisão do mundo real.

A principal contribuição deste trabalho é uma metodologia robusta que preenche essa lacuna. Foi desenvolvido e validado um modelo de simulação de eventos discretos de um terminal intermodal de grãos que integra, de forma explícita, métricas para a fadiga do operador e para a equidade na distribuição dos tempos de espera dos caminhoneiros. Essas métricas, fundamentadas na literatura de fatores humanos e teoria da justiça, foram incorporadas como funções-objetivo em um *framework* de otimização multi-objetivo, ao lado de indicadores operacionais tradicionais como vazão e utilização de ativos.

O resultado fundamental da aplicação do algoritmo NSGA-II a este modelo socio-técnico foi a geração de uma fronteira de Pareto que torna visíveis e quantificáveis os *trade-offs* inerentes entre eficiência, bem-estar e equidade. A análise demonstrou que é tecnicamente viável modelar e otimizar para esses objetivos concorrentes, movendo a discussão de um plano puramente qualitativo para uma análise de decisão baseada em dados. A metodologia fornece aos gestores um “menu de políticas”, onde cada opção tem seus custos e benefícios operacionais e éticos claramente delineados, permitindo uma tomada de decisão mais consciente e alinhada aos valores organizacionais.

Além disso, é imperativo reconhecer uma limitação ética estrutural do modelo proposto. O algoritmo de otimização, ainda que eficiente em explorar *trade-offs* entre desempenho operacional e métricas éticas internas (como fadiga e Coeficiente de Gini dos tempos de espera), permanece essencialmente cego para os impactos que ocorrem fora dos portões do terminal. Assim, nada impede que uma solução considerada “ótima” no escopo do modelo corresponda, na prática, a uma operação em ritmo frenético, sem pausas, que minimiza filas internas enquanto transfere todo o fardo da espera para o espaço público.

Nesse cenário, caminhões podem formar filas extensas na malha viária adjacente, gerando poluição do ar e sonora, congestionamento urbano, estresse e perda de tempo para motoristas, além de custos sociais difusos para a comunidade circundante. Este é um exemplo clássico de “transferência de custos” (*cost-shifting*), fenômeno conhecido na literatura de externalidades, mas eticamente invisível para o modelo quando não se incluem variáveis ou restrições que capturem essas consequências além dos limites físicos do sistema simulado.

Portanto, uma solução que apresente um Coeficiente de Gini interno mais baixo pode, paradoxalmente, ser socialmente indesejável no plano mais amplo. A variável `ConditionTruckArrivalMaxTrucks`, embora útil como restrição operacional, funciona apenas como um *proxy* fraco e insuficiente para conter esta dinâmica complexa de impactos deslocados.

Em suma, este trabalho oferece um caminho prático para a implementação da “ética desde o projeto” (*ethics by design*) em sistemas logísticos complexos. Ao transformar princípios éticos em objetivos de otimização mensuráveis, a metodologia aqui apresentada capacita as organizações a

projetar e operar sistemas que não são apenas eficientes, mas também mais justos e humanos. Este alinhamento entre inovação tecnológica e valores sociais não é apenas uma aspiração da Indústria 5.0, mas um requisito essencial para a sustentabilidade e a legitimidade das operações industriais na era digital.

VIII. LIMITAÇÕES E PESQUISAS FUTURAS

A demonstração de rigor acadêmico exige um reconhecimento transparente das limitações inerentes a qualquer estudo, bem como a identificação de caminhos promissores para a evolução da pesquisa. Esta seção cumpre esse papel, contextualizando os achados do presente trabalho e delineando uma agenda para investigações futuras.

A. Limitações do Estudo

Apesar da robustez da metodologia proposta, algumas limitações devem ser explicitamente declaradas:

- **Validação do Modelo de Fadiga:** Uma limitação crítica deste trabalho reside na ausência do componente circadiano (Processo C) no submodelo de fadiga. Essa omissão impede o modelo de diferenciar adequadamente o impacto da fadiga entre turnos diurnos e noturnos, o que é especialmente relevante para operações que funcionam em regime 24/7. A falta dessa dimensão limita significativamente o realismo fisiológico da simulação e, portanto, a capacidade do modelo de fornecer estimativas confiáveis de risco operacional ou impacto na segurança associados a diferentes janelas de trabalho. O submodelo de fadiga, embora conceitualmente fundamentado em modelos biomatemáticos, é uma abstração matemática. Seus parâmetros (`TaxaAcumulo`, `TaxaRecuperacao`) foram calibrados com base na lógica operacional e em dados de literatura, mas não foram validados empiricamente com dados fisiológicos (p. ex., EEG, variabilidade da frequência cardíaca) ou subjetivos (p. ex., Escala de Sonolência de Karolinska - KSS) coletados dos operadores reais do terminal. Além disso, a relação entre o `NivelFadiga` calculado pelo modelo e a probabilidade de erro humano não foi quantificada, embora a literatura de análise de confiabilidade humana (*Human Reliability Analysis* - HRA) sugira uma forte correlação, onde a fadiga é um Fator de Modelagem de Desempenho (*Performance Shaping Factor* - PSF) significativo [52], [71], [78].
- **Limites do Sistema e Externalidades:** Conforme discutido na Seção II-D, o modelo de simulação opera dentro de limites bem definidos, que se encerram nos portões do terminal. Esta delimitação, necessária para a tratabilidade computacional, exclui a análise de importantes externalidades negativas. O modelo não captura diretamente os custos sociais impostos à comunidade local, como o congestionamento de tráfego, a poluição do ar e a poluição sonora geradas pela fila de caminhões que se forma nas vias de acesso ao terminal [35]–[37]. A otimização de operações internas pode, paradoxalmente, agravar esses problemas externos.

- **Simplificação da Métrica de Equidade:** O conceito de equidade foi focado em um único grupo de stakeholders: os caminhoneiros, medido pela distribuição de seus tempos de espera. Uma análise mais abrangente poderia considerar outras dimensões da justiça, como a equidade na distribuição da carga de trabalho e, consequentemente, da fadiga entre os diferentes operadores e equipes. Esta dimensão de equidade interna não foi um objetivo explícito da otimização.

B. Direções para Pesquisas Futuras

As limitações identificadas não diminuem o valor do estudo, mas, ao contrário, abrem fronteiras claras para futuras pesquisas. As seguintes direções são propostas:

- **Validação e Enriquecimento do Modelo de Fadiga:** Como prioridade máxima para aprimorar o realismo fisiológico do modelo, recomenda-se o desenvolvimento e a incorporação de um módulo de ritmo circadiano (Processo C), permitindo capturar a interação entre o acúmulo de fadiga (Processo S) e as variações naturais de alerta ao longo das 24 horas. A implementação de tal módulo alinharia o modelo a práticas consagradas em indústrias críticas, onde o risco da fadiga noturna é fator determinante para planejamento de turnos, mitigação de erros humanos e políticas de segurança operacional. Uma extensão natural e valiosa deste trabalho seria a realização de estudos de campo para coletar dados fisiológicos e subjetivos dos operadores em um ambiente real. Tais dados permitiriam uma calibração e validação muito mais rigorosas do modelo de fadiga, aumentando seu poder preditivo e sua fidelidade à realidade. Além disso, o modelo poderia ser enriquecido para incluir a relação probabilística entre o nível de fadiga e a ocorrência de erros operacionais, permitindo a otimização direta da segurança.
- **Expansão dos Limites do Sistema para Incluir Externalidades:** Pesquisas futuras deveriam visar a expansão dos limites do sistema, integrando o modelo do terminal com um modelo de simulação de tráfego da rede viária circundante. Isso permitiria quantificar as externalidades (p. ex., tempo médio de bloqueio de vias, emissões de CO_2 da fila externa) e incluí-las como novas funções-objetivo na otimização, buscando soluções que sejam ótimas não apenas para o terminal, mas também para a comunidade.
- **Otimização Robusta e Adaptativa sob Incerteza Profunda:** O atual *framework* de otimização encontra soluções ótimas para um conjunto fixo de parâmetros operacionais. No entanto, o mundo real é caracterizado por incerteza profunda (*deep uncertainty*), como falhas inesperadas de equipamentos, picos de demanda não previstos, ou interrupções na cadeia de suprimentos. Uma direção de pesquisa avançada seria aplicar técnicas de Otimização Robusta ou Tomada de Decisão Robusta (*Robust Decision Making* - RDM). O objetivo não seria mais encontrar a solução "ótima" para um cenário esperado, mas sim identificar políticas operacionais "robustas" que tenham um bom desempenho em uma ampla gama de futuros plausíveis, mesmo que não sejam ótimas em nenhum deles.

- **Implementação como um Gêmeo Digital (*Digital Twin*) para Controle em Tempo Real:** A evolução final desta linha de pesquisa é a transição de um modelo de simulação para planejamento offline para um Gêmeo Digital operacional online. Tal sistema seria continuamente alimentado com dados em tempo real do terminal (p. ex., status dos sensores dos equipamentos, localização de veículos via GPS, tamanho real das filas). O *framework* de otimização executaria em ciclos curtos ou em resposta a eventos de perturbação (p. ex., a quebra de um tombador), gerando e recomendando agendamentos e políticas operacionais adaptativas em tempo real. Isso transformaria o modelo de uma ferramenta de análise estratégica em uma ferramenta de controle tático e resiliência operacional, representando um passo significativo em direção a operações logísticas verdadeiramente inteligentes e autônomas. Esta visão conecta diretamente o trabalho de mestrado a duas das áreas mais avançadas da engenharia de sistemas e da Indústria 4.0, demonstrando um roteiro claro para o impacto futuro da pesquisa.

REFERÊNCIAS

- [1] R. C. Saha and K. B. Khalil, "Policy Formulations to Establish More Dry Port Infrastructures to Increase Seaport Efficiency, Productivity, and Competitiveness in Bangladesh," *Future Transportation*, vol. 5, no. 2, p. 69, Jun. 2025, doi: 10.3390/futuretransp5020069.
- [2] P. Srisurin, P. Pimpanit, and P. Jarumaneeraj, "Evaluating the long-term operational performance of a large-scale inland terminal: A discrete event simulation-based modeling approach," *PLoS ONE*, vol. 17, no. 12, p. e0278649, Dec. 2022, doi: 10.1371/journal.pone.0278649.
- [3] L. Liu, J. Deng, and Y. Tang, "A Dynamic Management and Integration Framework for Models in Landslide Early Warning System," *ISPRS International Journal of Geo-Information*, vol. 12, no. 5, Art. no. 5, May 2023, doi: 10.3390/ijgi12050198.
- [4] H. Shi, R. Venkatesha Prasad, V. S. Rao, and I. G. M. M. Niemegeers, "A Fairness Model for Resource Allocation in Wireless Networks," in *NETWORKING 2012 Workshops*, Z. Becvar, R. Bestak, and L. Kencl, Eds., Berlin, Heidelberg: Springer, 2012, pp. 1–9. doi: 10.1007/978-3-642-30039-4_1.
- [5] V. Xinying Chen and J. N. Hooker, "A guide to formulating fairness in an optimization model," *Ann Oper Res*, pp. 1–39, Apr. 2023, doi: 10.1007/s10479-023-05264-y.
- [6] Sustainability, "AI Ethics Transportation," *Sustainability Directory*. Accessed: Jun. 08, 2025. [Online]. Available: <https://sustainability-directory.com/term/ai-ethics-transportation/>
- [7] T. Adewale, "(PDF) The Ethical and Security Challenges of AI Adoption in Logistics," *ResearchGate*. Accessed: Jun. 08, 2025. [Online]. Available: https://www.researchgate.net/publication/389987803_The_Ethical_and_Security_Challenges_of_AI_Adoption_in_Logistics
- [8] Sustainability, "Human-Centric Logistics." Accessed: Jun. 08, 2025. [Online]. Available: <https://energy.sustainability-directory.com/term/human-centric-logistics/>
- [9] Y. K. Tse, R. Dou, Z. Kiu, and P. Wang, "Human-centric AI in Operations and Supply Chain Management." Accessed: Jun. 08, 2025. [Online]. Available: https://think.taylorandfrancis.com/special_issues/human-centric-ai-oscm/
- [10] P. Radanliev, "AI Ethics: Integrating Transparency, Fairness, and Privacy in AI Development," *ResearchGate*, Feb. 2025, Accessed: Jun. 08, 2025. [Online]. Available: https://www.researchgate.net/publication/388803359_AI_Ethics_Integrating_Transparency_Fairness_and_Privacy_in_AI_Development
- [11] J. C. S. Lima, F. O. Silva and D. D. Ferreira, "A Urgência Ética da IA Centrada no Humano na Era da Automação do Trabalho," unpublished.
- [12] É. Marcon, M. Soliman, W. Gerstlberger, and A. G. Frank, "Sociotechnical factors and Industry 4.0: an integrative perspective for the adoption of smart manufacturing technologies," *JMTM*, vol. 33, no. 2, pp. 259–286, Feb. 2022, doi: 10.1108/JMTM-01-2021-0017.

- [13] B. Andres, M. Diaz-Madroñero, A. L. Soares, and R. Poler, "Enabling Technologies to Support Supply Chain Logistics 5.0,11 IEEE Access, vol. 12, pp. 43889–43906, 2024, doi: 10.1109/ACCESS.2024.3374194.
- [14] N. Dacre, Yan ,Jingyang, Frei ,Regina, Al-Mhdawi ,M. K. S., and H. and Dong, "Advancing sustainable manufacturing: a systematic exploration of Industry 5.0 supply chains for sustainability, human-centricity, and resilience," *Production Planning & Control*, 2024, pp. 1–30, doi: 10.1080/09537287.2024.2380361.
- [15] A. Kushe, S. Das, "Global Research and Analytics Firm — Investment, Business, Intellectual Property Research and Business Valuation Services," Aranca, 2023, Accessed: Jun. 08, 2025. [Online]. Available: <https://www.aranca.com/knowledge-library/articles/business-research/logistics-5.0-difference-with-a-human-touch>
- [16] J. C. S. Lima, F. O. Silva and D. D. Ferreira, "Ethical challenges in Artificial Intelligence: Navigating the boundaries of responsible innovation," unpublished.
- [17] M. A. Bazel, F. Mohammed, A. O. Baarimah, G. Alawi, A.-B. A. Al-Mekhlafi, and B. Almuhaaya, "The Era of Industry 5.0: An Overview of Technologies, Applications, and Challenges," in *Advances in Intelligent Computing Techniques and Applications*, F. Saeed, F. Mohammed, and Y. Fazea, Eds., Cham: Springer Nature Switzerland, 2024, pp. 274–284. doi: 10.1007/978-3-031-59707-7_24.
- [18] B. Avi-Itzhak and H. Levy, "On measuring fairness in queues," *Advances in Applied Probability*, vol. 36, no. 3, pp. 919–936, Sep. 2004, doi: 10.1239/aap/1093962217.
- [19] D. Raz, B. Avi-Itzhak, and H. Levy, "A resource-allocation queueing fairness measure: Properties and bounds," in *Proceedings of the ACM SIGMETRICS International Conference on Measurement and Modeling of Computer Systems*, 2004, pp. 130–141, doi: 10.1145/1012888.1005704.
- [20] A. Wierman and M. Harchol-Balter, "Classifying scheduling policies with respect to unfairness in an M/GI/1," in *Proceedings of ACM SIGMETRICS*, 2003, pp. 238–249, doi: 10.1145/781027.781055.
- [21] W. Whitt, "The amount of overtaking in a network of queues," *Networks*, vol. 14, no. 3, pp. 411–426, 1984, doi: 10.1002/net.3230140307.
- [22] W. Sandmann, "Analysis and optimization of fairness in multi-queue systems with server transitions," *Performance Evaluation*, vol. 59, no. 1, pp. 45–71, Jan. 2005, doi: 10.1016/j.peva.2004.06.005.
- [23] F. P. Kelly, A. K. Maulloo, and D. K. H. Tan, "Rate control for communication networks: shadow prices, proportional fairness and stability," *Journal of the Operational Research Society*, vol. 49, no. 3, pp. 237–252, Mar. 1998, doi: 10.1057/palgrave.jors.2600531.
- [24] T. Bonald, L. Massoulié, A. Proutière, and J. Virtamo, "A queueing analysis of max-min fairness, proportional fairness and balanced fairness," *Queueing Systems*, vol. 53, no. 1–2, pp. 65–84, May 2006, doi: 10.1007/s11134-006-6576-3.
- [25] Anylogic, "Grain Terminal - AnyLogic Cloud." Accessed: Jun. 08, 2025. [Online]. Available: <https://cloud.anylogic.com/model/7bfc5647-cd0f-4a66-ae42-22c866500da4?mode=SETTINGS&tab=GENERAL>
- [26] S. Rudenko, A. Shakhov, I. Lapkina, O. Shumylo, M. Malaksiano, and I. Horchynskyi, "Multicriteria Approach to Determining the Optimal Composition of Technical Means in the Design of Sea Grain Terminals," *Transactions on Maritime Science*, vol. 11, no. 1, Art. no. 1, Apr. 2022, doi: 10.7225/toms.v11.n01.003.
- [27] D. H. King and H. S. Harrison, "Open-source simulation software 'JaamSim,'" in 2013 Winter Simulations Conference (WSC), Dec. 2013, pp. 2163–2171. doi: 10.1109/WSC.2013.6721593.
- [28] JaamSim Development Team, "JaamSim Release Notes 2025-04," JaamSim. Accessed: Jun. 08, 2025. [Online]. Available: <https://jaamsim.com/docs/JaamSim%20Release%20Notes%202025-04.pdf>
- [29] H. King and H. S. Harrison, "JaamSim - Discrete-event Simulation Software," JaamSim- Discrete-event Simulation Software. Accessed: Jun. 08, 2025. [Online]. Available: <https://www.jaamsim.com/>
- [30] JaamSim Pro, "Mine-to-port supply chain model," (Oct. 08, 2022). Accessed: Jun. 08, 2025. [Online Video]. Available: <https://www.youtube.com/watch?v=A8mhfmYA-Vs>
- [31] K. Deb, A. Pratap, S. Agarwal, and T. Meyarivan, "A fast and elitist multiobjective genetic algorithm: NSGA-II," *IEEE Trans. Evol. Comput.*, vol. 6, no. 2, pp. 182–197, Apr. 2002, doi: 10.1109/4235.996017.
- [32] JaamSim software, "User Manual," JaamSim Software, Aug. 10, 2021. Accessed: Jun. 01, 2025. [Online]. Available: <https://jaamsim.com/downloads.html>
- [33] S. Lee, "System Boundaries in Practice," Number Analytics, 2025. Accessed: Jun. 22, 2025. [Online]. Available: https://www.numberanalytics.com/blog/system-boundaries-in-practice#google_vignette
- [34] M. Wayne, "Defining System Boundaries: Best Practices," *Requi Systems Engineering Articles*, 2024. Accessed: Jun. 22, 2025. [Online]. Available: <https://requi.io/articles/defining-system-boundaries-best-practices>
- [35] S. Lee, "A New Look at My Supply Chain Externalities," *Number Analytics*, Apr. 22, 2025. Accessed: Jun. 22, 2025. [Online]. Available: <https://www.numberanalytics.com/blog/supply-chain-externalities-guide>
- [36] J.-P. Rodrigue, "The concept of externalities," In *The Geography of Transport Systems*, 6th ed. London: Routledge, 2024. doi: 10.4324/9781003343196.
- [37] P. Rajsic, "The Ethics of Externalities," *Mises Institute*, Sep. 28, 2019. Accessed: Jun. 28, 2025. [Online]. Available: <https://mises.org/mises-daily/ethics-externalities>
- [38] B. Avi-Itzhak, H. Levy, and D. Raz, "Quantifying Fairness in Queuing Systems: Principles, Approaches, and Applicability," *Probability in the Engineering and Informational Sciences*, vol. 22, no. 4, pp. 495–517, Oct. 2008, doi: 10.1017/S0269964808000302.
- [39] Queue-it, "What Is a First-in-first-out (FIFO) Queue?," Queue-it. Accessed: Jun. 23, 2025. [Online]. Available: <https://queue-it.com/queue-first-in-first-out/>
- [40] M. Kim and S. Chai, "Impact of Justice in the Supply Chain Relationship on Implementing Supply Chain Integration," *International Journal of Supply Chain Management*, vol. 8, no. 6, 2019. Accessed: Jun. 01, 2025. [Online]. Available: <https://www.ojs.excelingtech.co.uk/index.php/IJSCM/article/viewFile/2950/2139>
- [41] C. Lee, S. Kim, and C. Kim, "Impact of Justice and Information Sharing on Logistics Performance in Supply Chain," *Journal of Distribution Science*, vol. 22, no. 8, pp. 137–145, 2024, doi: 10.15722/jds.22.08.202408.137.
- [42] A. R. Hofer, A. M. Knemeyer, and P. R. Murphy, "The Roles of Procedural and Distributive Justice in Logistics Outsourcing Relationships," *Journal of Business Logistics*, vol. 33, no. 3, pp. 196–209, Sep. 2012, doi: 10.1111/j.2158-1592.2012.01052.x.
- [43] D. B. McFarlin and P. D. Sweeney, "Research Notes. Distributive and Procedural Justice as Predictors of Satisfaction with Personal and Organizational Outcomes," *Academy of Management Journal*, vol. 35, no. 3, pp. 626–637, Aug. 1, 1992, doi: 10.5465/256489.
- [44] D. Raz, B. Avi-Itzhak, and H. Levy, "Locality of reference and the use of sojourn time variance for queue fairness," *School of Computer Science — Tel-Aviv University*, May 2005, Accessed: Jun. 22, 2025. [Online]. Available: <https://citeseerx.ist.psu.edu/document?repid=rep1&type=pdf&doi=cc5afbdbb29036ffb2879478a32931466a2bf551>
- [45] R. Dorfman, "A Formula for the Gini Coefficient," *The Review of Economics and Statistics*, vol. 61, no. 1, pp. 146–149, 1979, doi: 10.2307/1924845.
- [46] A. Hayes, S. Anderson, and D. Rubin, "Gini Index Explained and Gini Coefficients Around the World," *Investopedia*. Accessed: Jun. 24, 2025. [Online]. Available: <https://www.investopedia.com/terms/g/gini-index.asp>
- [47] Wikipedia, "Gini coefficient," Wikipedia. May 27, 2025. Accessed: Jun. 28, 2025. [Online]. Available: https://en.wikipedia.org/w/index.php?title=Gini_coefficient&oldid=1292519239
- [48] The Core Econ Team, "5.12 Measuring economic inequality: The Gini coefficient," In *The Economy 2.0*. Coreecon, 2024. Accessed: Jun. 28, 2025. [Online]. Available: <https://www.core-econ.org/the-economy/microeconomics/05-the-rules-of-the-game-12-measuring-economic-inequality.html>
- [49] R. M. Scism, "Human Fatigue in 24/7 Operations: Law Enforcement Considerations and Strategies for Improved Performance," *Police Chief Magazine*, 2017. Accessed: Jun. 18, 2025. [Online]. Available: <https://www.policechiefmagazine.org/human-fatigue-in-247-operations/>
- [50] Eddystone Team, "Human Factors – Managing fatigue in 24/7 operations," Eddystone, Mar. 30, 2017. Accessed: Jun. 18, 2025. [Online]. Available: <https://www.eddistone.com/human-factors-managing-fatigue-in-24-7-operations/>
- [51] S. E. Lerman et al., "Fatigue Risk Management in the Workplace," *Journal of Occupational & Environmental Medicine*, vol. 54, no. 2, pp. 231–258, Feb. 2012, doi: 10.1097/JOM.0b013e318247a3b0.
- [52] W. G. Bridges and W. Lane, "The Impact of Human Factors on LOPA (and other risk assessment methods)," *Process Improvement Institute*, 2013, Accessed: Jun. 05, 2025. [Online]. Available: https://process-improvement-institute.com/_downloads/Impact_of_Human_Error_on_LOPA_website.pdf

- [53] IATA, "Applications of Biomathematical Fatigue Models," In Human Factors Global Survey for Airlines, International Air Transport Association, 2025. Accessed: Jun. 18, 2025. [Online]. Available: <https://www.iata.org/en/programs/safety/operational-safety/fatigue-risk/>
- [54] C. 24/7 W. Solutions, "Measuring Worker Fatigue Risk: 4 Reasons Your Operation Should Be Using a Biomathematical Fatigue Model," 24/7 Global Workforce Solutions — CIRCADIAN® Technologies, Jul. 17, 2024. Accessed: Jun. 29, 2025. [Online]. Available: <https://circadian.com/blog/fatigue-risk>
- [55] S. M. Riedy, D. Fekedulegn, M. Andrew, B. Vila, D. Dawson, and J. Violanti, "Generalizability of a biomathematical model of fatigue's sleep predictions," *Chronobiol Int*, vol. 37, no. 4, pp. 564–572, Apr. 2020, doi: 10.1080/07420528.2020.1746798.
- [56] T. Raslear, "Criteria and procedures for validating biomathematical models of human performance and fatigue: Procedures for analysis of work schedules — FRA," Federal Railroad Administration, Jul. 19, 2013. Accessed: Jun. 05, 2025. [Online]. Available: <https://railroads.dot.gov/elibrary/criteria-and-procedures-validating-biomathematical-models-human-performance-and-fatigue>
- [57] M. M. Mallis, S. Mejdal, T. T. Nguyen, and D. F. Dinges, "Summary of the Key Features of Seven Biomathematical Models of Human Fatigue and Performance," *Aviation, Space, and Environmental Medicine*, vol. 75, no. 3, Mar. 2004. Accessed: Jun. 01, 2025. [Online]. Available: https://www.med.upenn.edu/uep/assets/user-content/documents/Mallis_etal_ASEM_75_3_2004.pdf
- [58] T. Åkerstedt and S. Folkard, "The Three-Process Model of Alertness and Its Extension to Performance, Sleep Latency, and Sleep Length," *Chronobiology International*, vol. 14, no. 2, pp. 115–123, Jan. 1997, doi: 10.3109/07420529709001149.
- [59] M. Ingre et al., "Validating and Extending the Three Process Model of Alertness in Airline Operations," *PLoS ONE*, vol. 9, no. 10, p. e108679, Oct. 2014, doi: 10.1371/journal.pone.0108679.
- [60] J. K. Devine, J. Choynowski, and S. R. Hursh, "Evaluation of a Biomathematical Modeling Software Tool for the Prediction of Risk in Flight Schedules Compared Against Incidence of Fatigue Reports," *Safety*, vol. 11, no. 1, p. 4, Jan. 2025, doi: 10.3390/safety11010004.
- [61] Z. Xianghan and L. Zhengwei, "Research on the Impact of Algorithmic Management on Employee Work Behavior in Platform Enterprises," *Journal of Economics and Management Sciences*, vol. 8, no. 2, p. p99, Apr. 2025, doi: 10.30560/jems.v8n2p99.
- [62] L. G. Mpedi and T. Marwala, "The Invisible Colleague — Balancing Algorithmic Management With Workers' Rights," *United Nations University*, Jan. 21, 2025. Accessed: Jun. 15, 2025. [Online]. Available: <https://unu.edu/article/invisible-colleague-balancing-algorithmic-management-workers-rights>
- [63] Haynes Law Firm, "The Growing Use of Algorithmic Management: Is Your Boss a Bot, and Is That Legal?," Haynes Law Firm, Jun. 6, 2025. Accessed: Jun. 21, 2025. [Online]. Available: <https://www.thehayneslawfirm.com/employment-info/employment-law-blog/the-growing-use-of-algorithmic-management-is-your-boss-a-bot-and-is-that-legal/>
- [64] M. D. Schlemmer and Z. W. Shine, "AI in the Workplace: The New Legal Landscape Facing US Employers," *Morgan, Lewis & Bockius LLP*, Jul. 1, 2024. Accessed: Jun. 29, 2025. [Online]. Available: <https://www.morganlewis.com/pubs/2024/07/ai-in-the-workplace-the-new-legal-landscape-facing-us-employers>
- [65] A. Bernhardt, L. Kresge, and R. Suleiman, "Data and Algorithms at Work: The Case for Worker Technology Rights," *UC Berkeley Labor Center*, Nov. 2021. Accessed: Jun. 12, 2025. [Online]. Available: <https://laborcenter.berkeley.edu/data-algorithms-at-work/>
- [66] X. Zhang, Y.-X. Liu, R. Wang, and E. M. Gutierrez-Farewik, "Multi-Objective-Based Human-in-the-Loop Optimization for Ankle Exoskeleton," *European Exoskeleton Conference*, Oct. 2023. Accessed: Mar. 12, 2025. [Online]. Available: https://www.aplusa-online.com/cgi-bin/md_aplusa/lib/all/lob/return_download.cgi/18_X_Zhang_Multi_Objective_Based_Human_in_the_loop_Optimization_for_Ankle_Exoskeleton.pdf?ticket=g_u_e_s_t&bid=6521&no_mime_type=0
- [67] CASA, "Biomathematical Fatigue Models Guidance Document Summary," *Australian Civil Aviation Safety Authority*, 2014. Accessed: Jun. 18, 2025. [Online]. Available: <https://www.iata.org/contentassets/5f976bb3ca2446f3a40e88b18dd61fb/b/condensed-version-of-casa-biomathematical-models-doc.pdf>
- [68] F. H. Cardoso, "Simulação de um modelo padrão para terminais integradores de grãos," *Universidade Federal de Minas Gerais*, 2014, [Online]. Available: <http://hdl.handle.net/1843/VRNS-9M8R8W>
- [69] Office of Railroad Policy and Development, "Procedures for Validation and Calibration of Human Fatigue Models: The Fatigue Audit InterDyne Tool," *Federal Railroad Administration*, Washington, DC, DOT/FRA/ORD-10/14, Oct. 2010. Accessed: Jun. 26, 2025. [Online]. Available: <https://www.phmsa.dot.gov/sites/phmsa.dot.gov/files/docs/technical-resources/pipeline/control-room-management/69061/trproceduresorvalidationandcalibrationfinal.pdf>
- [70] M. E. McCauley et al., "Fatigue risk management based on self-reported fatigue: Expanding a biomathematical model of fatigue-related performance deficits to also predict subjective sleepiness," *Transp Res Part F Traffic Psychol Behav*, vol. 79, pp. 94–106, May 2021, doi: 10.1016/j.trf.2021.04.006.
- [71] Wikipedia, "Human reliability," *Wikipedia*. Dec. 04, 2024. Accessed: Jun. 26, 2025. [Online]. Available: https://en.wikipedia.org/w/index.php?title=Human_reliability&oldid=1261189548
- [72] J. Robertson, T. Schmidt, F. Hutter, and N. Awad, "A Human-in-the-Loop Fairness-Aware Model Selection Framework for Complex Fairness Objective Landscapes," *AIES*, vol. 7, pp. 1231–1242, Oct. 2024, doi: 10.1609/aies.v7i1.31719.
- [73] L. Thiele, K. Miettinen, P. J. Korhonen, and J. Molina, "A Preference-Based Evolutionary Algorithm for Multi-Objective Optimization," *Evolutionary Computation*, vol. 17, no. 3, pp. 411–436, Sep. 2009, doi: 10.1162/evco.2009.17.3.411.
- [74] P.-L. Rodrigue, "Goodhart's Law: The Hidden Risk in Software Engineering Metrics," *Axify*, Apr. 1, 2025. Accessed: Jun. 28, 2025. [Online]. Available: <https://axify.io/blog/goodhart-law>
- [75] J. Z. Muller, "The Tyranny of Metrics," *Princeton University Press*, 2019. ISBN: 9781400889433.
- [76] É. Marcon, M. Soliman, W. Gerstlberger, and A. G. Frank, "Sociotechnical factors and Industry 4.0: an integrative perspective for the adoption of smart manufacturing technologies," *Journal of Manufacturing Technology Management*, vol. 33, no. 2, pp. 259–286, Sep. 2021, doi: 10.1108/JMTM-01-2021-0017.
- [77] A. Sayem, P. K. Biswas, M. M. A. Khan, L. Romoli, and M. Dalle Mura, "Critical Barriers to Industry 4.0 Adoption in Manufacturing Organizations and Their Mitigation Strategies," *Journal of Manufacturing and Materials Processing*, vol. 6, no. 6, Art. no. 6, Dec. 2022, doi: 10.3390/jmmp6060136.
- [78] M. Rasmussen and K. Laumann, "The evaluation of fatigue as a performance shaping factor in the Petro-HRA method," *Reliability Engineering and System Safety*, vol. 194, no. C, 2020. Accessed: Jun. 29, 2025. [Online]. Available: <https://ideas.repec.org/a/eee/reensys/v194y2020ics0951832017310402.html>

3 CONSIDERAÇÕES FINAIS

Atualmente, as deficiências estruturais da logística brasileira representam um desafio significativo para o agronegócio, especialmente no que se refere à exportação de grãos. Esses problemas, que vão desde a insuficiência de infraestrutura até a ineficiência no escoamento das safras, afetam diretamente a competitividade do setor no mercado internacional. No entanto, em um ambiente tão dinâmico e competitivo como o agronegócio, a capacidade de tomar decisões rápidas e bem fundamentadas se torna crucial para superar essas barreiras.

Nesse contexto, o uso de modelos computacionais avançados, como simulações e algoritmos de otimização, emerge como uma alternativa estratégica para enfrentar os entraves logísticos. Essas ferramentas permitem que gestores e operadores antecipem problemas, ajustem processos em tempo real e adotem estratégias mais eficientes. Assim, os modelos computacionais não apenas ajudam a mitigar os impactos das limitações estruturais, mas também proporcionam ganhos operacionais decisivos para manter a competitividade do Brasil como um dos maiores exportadores de grãos do mundo.

Partindo dessa premissa, este estudo investigou a aplicação de técnicas de inteligência artificial e simulação para a otimização de operações logísticas em terminais intermodais de grãos. A partir de uma fundamentação teórica estruturada em cinco eixos — otimização multiobjetivo com NSGA-II, simulação de eventos discretos, HCAI, confiança e aspectos psicossociais em IA, e modelos híbridos aplicados à logística — foi possível desenvolver uma abordagem metodológica sólida e coerente com os desafios operacionais, éticos e ambientais dos terminais logísticos brasileiros.

Dentre as principais conquistas deste estudo, destacam-se:

- a) a modelagem realista de um terminal multimodal usando simulação de eventos discretos, validada com base em dados operacionais;
- b) a implementação de um processo de otimização multiobjetivo com NSGA-II, capaz de identificar soluções não dominadas considerando critérios operacionais (vazão, utilização de ativos), sociais (fadiga, desigualdade distributiva) e ambientais (tempo de permanência como *proxy* da pegada de carbono);
- c) a integração explícita de considerações éticas e humanocêntricas na avaliação de cenários otimizados, promovendo a justiça distributiva e a equidade no uso de recursos.

A análise dos resultados permitiu identificar *trade-offs* claros entre os objetivos definidos, demonstrando que ganhos em produtividade frequentemente exigem sacrifícios em termos de equidade ou bem-estar operacional. Tal constatação reforça a importância de adotar abordagens de otimização que considerem múltiplos valores, não apenas métricas de eficiência.

No entanto, algumas limitações devem ser reconhecidas. A modelagem baseou-se em suposições simplificadoras, como perfis médios de operadores e parâmetros estáticos de chegada de caminhões. A variabilidade comportamental dos agentes humanos, assim como eventos disruptivos (clima, falhas técnicas), foram parcialmente abstraídos. Além disso, os critérios de avaliação ética e ergonômica poderiam ser refinados por meio de métodos empíricos de coleta de dados qualitativos, como entrevistas ou questionários com operadores reais.

A relação entre teoria e prática foi fortalecida ao conectar os fundamentos da simulação e da IA com a aplicação em um caso concreto de terminal logístico, promovendo um diálogo entre engenharia de produção, ciência da computação e ética aplicada. A literatura sobre Indústria 5.0, governança algorítmica e modelos híbridos encontrou, neste trabalho, um terreno fértil para validação interdisciplinar e aplicação prática.

Para estudos futuros, propõe-se:

- a) explorar modelos com agentes inteligentes adaptativos, que considerem aprendizagem e decisões locais dos caminhões e operadores;
- b) incorporar métricas mais robustas de impacto ambiental, como consumo energético e emissões de CO_2 reais, por meio de inventários e sensores;
- c) realizar estudos participativos com *stakeholders* do setor logístico para avaliar, qualitativamente, a aceitabilidade das soluções propostas;
- d) aplicar a abordagem a outros tipos de terminais (portuários, ferroviários, aeroportuários) e produtos (fertilizantes, *containers*).

Conclui-se, portanto, que a integração entre simulação, inteligência artificial e princípios éticos representa um caminho promissor para enfrentar os desafios complexos da logística moderna, permitindo não apenas ganhos em eficiência, mas também avanços significativos em sustentabilidade e justiça operacional.

REFERÊNCIAS

- ABAY, Y. *et al.* A discrete-event simulation study of multi-objective sales and operation planning: A case of ethiopian automotive industry. **IFAC-PapersOnLine**, v. 56, n. 2, p. 7826–7833, 2023. ISSN 2405-8963. 22nd IFAC World Congress. Disponível em: <https://www.sciencedirect.com/science/article/pii/S2405896323015501>.
- ALAHMED, Y. *et al.* Bridging the Gap Between Ethical AI Implementations. **International Journal of Membrane Science and Technology**, v. 10, n. 3, p. 3034–3046, out. 2023. ISSN 2410-1869. Publisher: Cosmos Scholars Publishing House. Disponível em: <https://cosmoscholars.com/phms/index.php/ijmst/article/view/2953>.
- ALMEDER, C.; PREUSSER, M. A hybrid optimization and simulation approach for supply chains. In: **Proceedings of EUROSIM 2007**. Ljubljana, Slovenia: [s.n.], 2007.
- ALVES, L. *et al.* Editorial: Human-centered artificial intelligence in industry 5.0. **Frontiers in Artificial Intelligence**, 2024.
- ANTONS, O.; ARLINGHAUS, J. C. Data-driven and autonomous manufacturing control in cyber-physical production systems. **Computers in Industry**, Elsevier, v. 141, p. 103711, 10 2022. ISSN 0166-3615.
- BOTTANI, E.; CASELLA, G. Discrete-event simulation in logistics and supply chain management: a scientometric perspective. **Production & Manufacturing Research**, 2024. 127papers analisados entre 2009–2022.
- BRASIL, M. d. T. **PLANO NACIONAL DE LOGÍSTICA PNL - 2025 Relatório Executivo**. 2018. Disponível em: <https://www.gov.br/transportes/pt-br/assuntos/planejamento-integrado-de-transportes/politica-e-planejamento/publicacoes/pnl2025.pdf>. Acesso em: 11 ago 2024.
- BREQUE, M.; NUL, L. D.; PETRIDIS, A. **Industry 5.0 - Towards a sustainable, human-centric and resilient European industry**. 2021. Disponível em: https://research-and-innovation.ec.europa.eu/knowledge-publications-tools-and-data/publications/all-publications/industry-50-towards-sustainable-human-centric-and-resilient-european-industry_en.
- CHOI, H.; LEE, S. Determinants of trust in artificial intelligence: A meta-analytic review. **Psychological Bulletin**, v. 149, n. 11, p. 1303–1327, 2023.
- CONAB. **Acompanhamento da safra brasileira - Grãos**. 2024. Disponível em: https://www.conab.gov.br/info-agro/safras/graos/boletim-da-safra-de-graos/item/download/56050_f25a1aab663edebd1c2dfc7aecb0f7fb. Acesso em: 06 fev 2025.
- DEB, K. *et al.* A fast and elitist multiobjective genetic algorithm: Nsga-ii. **IEEE transactions on evolutionary computation**, IEEE, v. 6, n. 2, p. 182–197, 2002.
- DIAZ-RODRIGUEZ, N. *et al.* Connecting the dots in trustworthy artificial intelligence: From ai principles, ethics, and key requirements to responsible ai systems and regulation. **Information Fusion**, 5 2023. Disponível em: <http://arxiv.org/abs/2305.02231>.

- DU, J. *et al.* Enhancing nsga-ii algorithm through hybrid strategy for optimizing maize water and fertilizer irrigation simulation. **Symmetry**, v. 16, n. 8, p. 1062, 2024.
- HUANG, C. *et al.* An overview of artificial intelligence ethics. **IEEE Transactions on Artificial Intelligence**, Institute of Electrical and Electronics Engineers Inc., 8 2022. ISSN 26914581.
- KPMG; MELBOURNE, U. of. **Trust, Attitudes and Use of Artificial Intelligence: A Global Study**. 2025. Relatório técnico. Pesquisa global com trabalhadores sobre confiança em IA.
- LIMA, J. C. S.; FERREIRA, D. D.; SILVA, F. O. Otimização socio-técnica de terminais intermodais: Uma abordagem de simulação multi-objetivo para equilibrar eficiência, bem-estar do trabalhador e equidade. **Manuscrito em revisão para publicação**, 2025. Artigo em revisão – Part II, Artigo 3 desta dissertação.
- LIMA, J. C. S.; SILVA, F. O.; FERREIRA, D. D. Ethical challenges in artificial intelligence: Navigating the boundaries of responsible innovation. **Manuscrito submetido para publicação**, 2025. Artigo submetido – Part II, Artigo 1 desta dissertação.
- LIMA, J. C. S.; SILVA, F. O.; FERREIRA, D. D. A urgência Ética da ia centrada no humano na era da automação do trabalho. **Manuscrito submetido para publicação**, 2025. Artigo submetido – Part II, Artigo 2 desta dissertação.
- MAHMOODI, A. *et al.* An enhanced nsga-ii algorithm with variable neighborhood search for multi-objective scheduling problems. **Expert Systems with Applications**, Elsevier, v. 215, p. 120375, 2025.
- MAYER, S.; CLASSEN, T.; ENDISCH, C. Modular production control using deep reinforcement learning: proximal policy optimization. **Journal of Intelligent Manufacturing**, Springer, v. 32, p. 2335–2351, 12 2021. ISSN 15728145. Disponível em: <https://link-springer-com.ez26.periodicos.capes.gov.br/article/10.1007/s10845-021-01778-z>.
- NEAL, E. *et al.* Transparent ai systems and their impact on operator trust and performance in logistics. **International Journal of Human–Computer Studies**, v. 170, p. 102748, 2024.
- NEUMANN, F. *et al.* Automation and worker well-being: Field survey in port and distribution center workers. **Journal of Industrial Psychology**, 2025. Estudo de campo em ambientes logísticos.
- OLUYISOLA, O. E.; SGARBOSSA, F.; STRANDHAGEN, J. O. Smart production planning and control: Concept, use-cases and sustainability implications. **Sustainability**, Multidisciplinary Digital Publishing Institute, v. 12, p. 3791, 5 2020. ISSN 2071-1050. Disponível em: <https://www.mdpi.com/2071-1050/12/9/3791/htm><https://www.mdpi.com/2071-1050/12/9/3791>.
- PAGOTTO, L. *et al.* Dynamic simulation budget allocation for hybrid simulation-optimization in logistics systems. **Simulation Modelling Practice and Theory**, 2025. Artigo em preparação.
- PASSALACQUA, M. *et al.* Human-centred ai in industry5.0: A systematic review. **International Journal of Production Research**, 2024. Revisão sistemática com 67 estudos.

- PERA, T. **Pesquisa quantifica perdas logísticas de soja e milho no Brasil**. 2017. Disponível em: <https://jornal.usp.br/ciencias/ciencias-agrarias/pesquisa-quantifica-perdas-logisticas-de-soja-e-milho-no-brasil>. Acesso em: 24 jul 2023.
- PEREIRA, R. d. S. **O setor portuário de Sergipe e Alagoas: fluxos de mercadorias e desenvolvimento regional**. Dissertação (Mestrado) — Universidade Federal de Uberlândia, 2020.
- PORTOCARRERO, M.; SILVA, R. Gargalos logísticos nos terminais intermodais de grãos no Brasil: Um estudo de caso. **Revista de Logística e Transporte**, 2021.
- RODRÍGUEZ-ESPINOSA, D. F. *et al.* Nsga-ii simheuristic to solve a multi-objective flexible flow shop problem under stochastic machine breakdowns. **Journal of Project Management (Canada)**, v. 9, n. 4, p. 493–512, 2024.
- ROŽANEC, J. M. *et al.* Human-centric artificial intelligence architecture for industry5.0 applications. **International Journal of Production Research**, v. 61, n. 20, p. 6847–6872, 2022.
- SIAU, K.; WANG, Z. Building trust in artificial intelligence, machine learning, and robotics. **Cutter Business Technology Journal**, v. 31, n. 2, p. 47–53, 2018.
- SILVA, C. D.; HALLOLUWA, T.; VYAS, D. A multi-layered research framework for human-centered ai: Defining the path to explainability and trust. **arXiv preprint**, arXiv:2504.13926, 2025.
- SOUZA, F. A. F. de. **Elaboração de um modelo de localização de cargas unitizadas agroindustriais em pátios portuários: Aplicação ao caso do terminal Portuário do Pecém**. Dissertação (Mestrado) — Universidade Federal do Ceará, 2002. Disponível em: <https://repositorio.ufc.br/handle/riufc/4870>.
- SRISURIN, P.; PIMPANIT, P.; JARUMANEEROJ, P. Evaluating the long-term operational performance of a large-scale inland terminal: A discrete event simulation-based modeling approach. **PLoS One**, v. 17, n. 12, p. e0278649, 2022.
- TEIXEIRA, P. E. F. **Desempenho de terminais hidroviários do corredor logístico Centro-Oeste: um estudo de multi-casos**. Dissertação (Mestrado) — Universidade Federal de Mato Grosso do Sul, 2010.
- THENARASU, M. *et al.* Multi-criteria scheduling of realistic flexible job shop: a novel approach for integrating simulation modelling and multi-criteria decision making. **arXiv preprint**, 2023. ArXiv:2308.03379.
- TOORAJIPOUR, R. *et al.* Artificial intelligence in supply chain management: A systematic literature review. **Journal of Business Research**, Elsevier, v. 122, p. 502–517, 1 2021. ISSN 0148-2963.
- VASSOS, G. *et al.* A simulation framework of procurement operations in the container logistics industry. **arXiv preprint**, 2023. ArXiv:2303.12765.
- VYHMEISTER, E.; CASTANE, G. G. When industry meets trustworthy ai: A systematic review of ai for industry 5.0. **arXiv preprint**, 2024. ArXiv:2403.03061.

WALTERSMANN, L. *et al.* Artificial intelligence applications for increasing resource efficiency in manufacturing companies—a comprehensive review. **Sustainability**, Multidisciplinary Digital Publishing Institute, v. 13, p. 6689, 6 2021. ISSN 2071-1050. Disponível em: <https://www.mdpi.com/2071-1050/13/12/6689/html><https://www.mdpi.com/2071-1050/13/12/6689>.

WOLSCHICK, D. *et al.* A comparative study of nsga-ii, nsga-iii and moea/d on many-objective optimization problems. **Applied Soft Computing**, Elsevier, v. 136, p. 110224, 2024.

XU, X. *et al.* Industry 4.0 and Industry 5.0—Inception, conception and perception. **Journal of Manufacturing Systems**, v. 61, p. 530–535, out. 2021. ISSN 02786125. Disponível em: <https://linkinghub.elsevier.com/retrieve/pii/S0278612521002119>.

YANG, S. *et al.* Verification of intelligent scheduling based on deep reinforcement learning for distributed workshops via discrete event simulation. **Advances in Production Engineering Management**, v. 17, p. 401–412, 12 2022. ISSN 18546250.

ZHOU, A. *et al.* Multiobjective evolutionary algorithms: A survey of the state of the art. **Swarm and Evolutionary Computation**, Elsevier, v. 1, n. 1, p. 32–49, 2011.

ÖZKUL, F.; SUTHERLAND, R.; WENZEL, S. Simulation-enhanced action-oriented process mining in production and logistics. In: UNIVERSITY OF THE BUNDESWEHR MUNICH. **Proceedings of the 2024 Simulation Technology Symposium**. [S.l.], 2024.