

Ethical Implications of AI-Driven Recruitment: A Multi-Perspective Study on Bias and Transparency in Digital Hiring Platforms

Fadlul musrifah*¹, Ika Ariani Hasanah²

Email: muser110@gmail.com; ika.ariani@gmail.com

¹Universitas Islam Negeri Walisongo, Semarang, Indonesia

²Universitas Islam Negeri Sunan Kalijaga, Yogyakarta, Indonesia

*Corresponding Author

Abstract

Integrating Artificial Intelligence (AI) into digital recruitment platforms has introduced significant enhancements in efficiency and decision-making, alongside complex ethical challenges regarding fairness, transparency, and accountability in candidate evaluation. This study investigates how leading AI-driven recruitment platforms articulate and operationalize ethical principles and whether these commitments are effectively translated into practice. Employing a qualitative exploratory design, the research analyzes official white papers, privacy policies, and AI ethics statements from LinkedIn, HireVue, Pymetrics, and ModernHire. Data was examined using AI-assisted text mining and thematic content analysis to identify ethical discourse patterns and assess the depth of implementation. The findings indicate that moral terms such as “fairness” and “bias” are cited frequently, with LinkedIn referencing them 27 times and HireVue 19 times. A comparative transparency assessment yielded scores of 8.5 out of 10 for LinkedIn, 7.2 for HireVue, 6.8 for Pymetrics, and 4.3 for ModernHire, while formal mechanisms for candidate appeals were absent on most platforms. This study contributes to the field by revealing a persistent gap between stated ethical ideals and operational practices in AI recruitment and by recommending the adoption of explainable AI, transparent auditing frameworks, and international regulatory standards. Such measures are essential to foster more accountable, equitable, and humane AI-based hiring processes.

Keywords: Algorithmic Bias, AI Recruitment Ethics, Transparency in Hiring, Explainable AI, Digital Hiring Platforms

Received on March 2025; Revised on March 2025; Accepted on April 2025; Published on April 2025.

I. INTRODUCTION

The adoption of AI within different industries has led to increased productivity, scalability, and decision-making precision. In human resource management, recruitment has been automated with the use of AI, which has significantly improved the candidate-to-job suitability rate and reduced the time required to fill positions. AI technologies are now integrated into platforms such as LinkedIn and HireVue, where algorithms evaluate resumes, video interviews, and rank the candidates according to how well they would fit the role. While these advancements have potential operational efficiencies, they come with deep ethical dilemmas such as fairness, responsibility, and openness. The discrimination based on AI algorithms is said to pose a risk if they act as gatekeepers to employment opportunities, due to schisms of gender, ethnic, or age bias that borrowing from deep-rooted training data perpetuates. Issues of discrimination are exacerbated by the fact that many AI models remain obscure to the candidate, and so do the standards they need to meet to be considered for such positions. The reality of the workplace

requires us not only to examine the efficiency with which algorithms do their tasks, but also how moral and just the results are.

Several studies have begun to document the ethical implications of AI in recruitment. For instance, found that AI-driven decision systems often reflect historical biases embedded in training datasets, which poses significant risks of reinforcing structural discrimination. Similarly, emphasize that while machine learning algorithms can improve efficiency in candidate selection, most platforms fail to provide transparency regarding how evaluations are conducted and what features influence final decisions, making the recruitment process opaque for applicants. Outlines principles such as inclusiveness and accountability, yet admits that operationalizing these ideals remains a challenge. Likewise, you should emphasize that AI systems trained on unbalanced or homogeneous datasets can unintentionally encode demographic and behavioral biases, thus making unfair assessments of candidates whose profiles deviate from dominant norms. Despite this insight, much of corporate disclosure on algorithm logic and ethical governance remains platitudinal or simply symbolic. What these patterns show is that while the topic has evoked academic and public interest, actual applications in the field where ethical safeguards take form with such depth and consistency are regrettably lacking.

There have been several studies initiated to document the ethical ramifications of AI in recruitment. For example, (Munifah et al., 2024) established that decision systems involving AI usually encapsulate historical biases in training data, with high risks for reinforcement of structural discrimination. In like manner, (Wahyuning & Sudibyo, 2024) point out that machine learning algorithms have the potential to increase efficiency in candidate choice, yet most platforms do not reveal how or on what basis evaluations take place and what characteristics determine conclusions, rendering the recruitment process unclear to candidates. (Nannini et al., 2024) delineates ideals such as accountability and inclusiveness, though it concedes that these ideals' operationalization remains an issue. In like manner, (Sumarlin & Kusumajaya, 2024) note that AI systems based on skewed or homogenous datasets have a high likelihood to embed unintentional demographic or behavioral biases, resulting in biased ratings of candidates whose profiles diverge from dominant norms. Despite these studies, corporate disclosures regarding algorithmic logic and ethical regulation are still largely on the surface or symbolic. Such trends indicate that, though the subject has attracted academic and popular attention, pragmatic applications of ethical protection in AI recruitment systems are still shallow and erratic.

While there has been a great leap forward in existing research on how operative efficiency and predictability in recruitment are carried out by AI, its ethical aspects remain underexplored. Research like that conducted by (Tavares & Ferrara, 2023) and (Ricart et al., 2022) has

demonstrated how AI is capable of amplifying current societal biases through historical datasets, but these pieces of research tend to fall short by not exploring institutional machinery facilitating such bias over digital recruitment platforms. (Almeida et al., 2025) have insisted on embedding ethical values in governing AI, but empirical evidence on how employment platforms operationalize such values in practice is still rare. Besides, corporate disclosures regarding ethical AI remain more high-sounding than high-impact. For instance, (Eke & Stahl, 2024) reported that most platforms provide ethical principles without operational definition or mechanisms for implementing them. (Thomaidou & Limniotis, 2025) also emphasize how algorithmic evaluations are transparent and how there are no mechanisms for redress for job applicants, a concern that is rarely studied from multi-actor or comparative perspectives. It is with this gap that this research aims to address by exploring how algorithmic bias is institutionalized and how transparency is (not) operationalized on top digital recruitment platforms.

The research is intended to pinpoint and examine with critical scrutiny major ethical issues arising within recruitment tools powered by AI, focusing on algorithmic bias as well as transparency deficiency underlying candidate screening practices. By exploring several different online employment platforms, this research will investigate how organizations perceive, manage, or avoid addressing discriminatory biases remaining in computer-assisted screening tools. A key research question informing this inquiry is: How do online employment platforms operationalize fairness and transparency within their AI-powered recruitment tools? Moreover, this research examines how ethical obligations outlined in company white papers are manifested in actual platform design and decision-making practices. By taking a multi-lens examination based on analysis of white papers, privacy policies, and AI ethics declarations, this research aspires to uncover systemic gaps and recommend more accountable and fairer frameworks of AI governance in recruitment. Ultimately, findings are meant to feed into academic scholarship as well as policy proposals for developing transparent, bias-sensitivity-aware, and ethics-reliant AI-based recruitment practices.

II. LITERATURE REVIEW

A. Theoretical Framework

1. Algorithmic Bias and Black Box AI in Recruitment Decision-Making

Thus, algorithmic bias within recruitment tools is an issue under a speedily advancing pace, which needs to be investigated most seriously under the broad stroke debate over AI ethics. Machine learning algorithms replicate trends of discrimination embedded in decisions of the past because of training from past employment data, as stated by (Pagano et al., 2023). Such biases get unintentionally programmed into algorithms driving the procedures of screening, shortlisting, or

ranking a candidate. The issue becomes very pronounced when data consists of antecedent discrepancies on grounds like gender, race, or even socioeconomic background. Instead of having a socio-political system built by humans without any significant scale, algorithmic systems end up amplifying and propelling their biases. This poses a dilemma of such complexity that it gives rise to incompatible organizational efforts to identify the unique variable or association that results in an unequal outcome in the recruitment process.

The "black box," or rather, the opaque nature of AI decisions, hampers transparency and accountability in the recruitment process. According to (Hassija et al., 2024), there is a black-box system which is not easily interpreted either by the users of the system or the people whose lives the output of the system influences. This has got to do with the fact that, on the one hand, the technical complexity of the models is such and on the second point, the AI providers impose many proprietary restrictions because such automatic suggestions are used by human resources professionals without any sight of the logic or attributes involved in making such inferences. This is a limitation for feedback and reduces opportunities for review. There is no way to see the internal operations, and thus, this is a barrier to unearthing biases behind certain patterns that disfavor some segments of job applicants.

Based on a thorough and careful examination of fairness, the challenges of addressing the ethical design versus operational objectives framework also emerge. As suggested by (Chen et al., 2023), the potential for fairness constraints to be introduced into the design of algorithms must also accommodate competing priorities of the organisation, such as performing efficiently or having data available. The inferred fairness of an attribute could be reinstated by using a seemingly neutral input that is capable of functioning as an indirect proxy for an attribute that falls under scrutiny through its covert associations. In this manner, factors like educational achievements, location, or even gaps in employment might correlate with broader social patterns of discrimination and disparity. Such underlying relationships between data attributes and historical disadvantages add to the complexities for those attempting to build systems that would yield unbiased conclusions to fulfill the requirements put forward by recruitment entities.

The missing auditing frameworks pose an additional layer of complexity to ensuring fair functioning in AI-based recruitment systems. (Felländer et al., 2022) highlight that most organizations do not use standardized protocols to examine the ethical effects of AI tools before use in real-life environments. Internal assessments are usually narrow and shallow, while outside audits are exceptional, considering they are constrained by poor access to proprietary details. When platforms dictate availability regarding performance information and algorithmic rationale, then outside players are more likely to struggle with determining fairness in the final technology.

Besides, recruitment algorithms get updated or retrained, making ethical performance consistency over a long-term span difficult to keep track of. These details demonstrate the necessity for ongoing scrutiny of both technical protocols and governing arrangements in AI-based recruitment systems.

2. Ethical Principles in Digital Recruitment: Fairness, Transparency, and Explainability

This is impartial and fair treatment of job applicants of various social and diverse demographics. As defined by (Lamers et al., 2022), it is a contextually dependent phenomenon, i.e., on inherent characteristics and operational features of the algorithm alongside standard rules and laws in society. Fairness in algorithmic recruitment systems is generally established by the extent to which results are similar for individuals manifesting in identical qualifications, as well as for individuals disfavoured historically by the system of algorithmic recruitment. Most of their unfairness roots from biased data that have historical roots and are still manifest in modern recruitment. Finding fairness in digital recruitment involves analyzing how ethics are integrated into data that speaks to society dynamics, and algorithmic development reflects specific values about decision-making and merit in the algorithmic development process.

Traditionally, AI-enabled recruitment transparency is understood concerning how transparent the algorithmic procedures and decision criteria will be for stakeholders. According to (Balasubramaniam et al., 2023), however, actual transparency means much more than describing an algorithm in technical terms; it means users' interpretive competence to make sense of these systems and use them effectively. While most e-recruiting sites provide superficial disclosures to vague mentions of input parameters or scoring parameters, such disclosures would never empower users to evaluate system actions authoritatively. The constitution of transparency as an interpretive property establishes a more differentiated perspective, where ethical sufficiency in disclosure hinges on users' capabilities to interpret provided information and use it to question or challenge algorithmic actions.

Explainability in an AI context is defined as the ability of a system to provide reasons behind its outputs in terms accessible to human stakeholders, particularly in consequential decision-making situations like employment. (Tursunaliyeva et al., 2024) highlights the intrinsic trade-off at play here, whereby systems tuned for high predictive accuracy tend to become uninterpretable in terms of how they make decisions. In recruitment scenarios, this untransparency can make outcomes untraceable and erode trust, as is especially the case when models utilize high numbers of variables and non-linear functions. The use of post hoc explanation methods is likely in attempts to overcome this limitation, though these can provide simplified explanations and may or may not correspond to how the system works internally. The existence of such interpretive lacunae serves

as a reminder that explainability is not just about technical measures, but also communicative intelligibility and trust on the part of users.

Such ethical requirements as fairness, transparency, and explainability ought to be perceived not merely as technical attributes but as outcomes of governance processes driven by institutional norms and normative priorities. The authors (Ortega-Bolaños et al., 2024) argue that ethical considerations should be consistently considered throughout the life cycle of AI system development—from data selection and model specification to deployment and user interaction. Decisions about what constitutes valid data, what compromises are tolerable, and what interests to organize are usually made within organizations that might be implicitly fostering the goal of efficiency or profit-making premise rather than that of social accountability. Without systematic institutional frameworks that scrutinize and ensure ethical practice, AI-based systems in employment may only have variable or incomplete commitment to ethical integrity, subject to organisational and societal constraints.

B. Previous Studies

1. Studies on Gender and Racial Bias in AI Recruitment Systems

Gender and racial distinctions in algorithm-based recruitment measures have been widely debated in academic research addressing ethics in AI. (Rigotti & Fosch-Villaronga, 2024) note that most AI-based recruitment tools, while created to make recruitment more efficient, tend to reproduce prevailing disparities based on the data they are trained on. The tools often use proxies, including educational attainment, past employment, or geographic area, which may have an association with race or gender and contribute to amplifying systemic biases. The intrinsic nature of these patterns gets replicated and affects candidate ordering as well as access to job opportunities for marginalized groups. Automation may provide a veneer of impartiality, but outputs produced by these systems will likely replicate prevailing social dynamics stemming from previous human decision-making. The examination uncovers how technical aspects in AI recruitment tools are deeply entangled with historical information and normative judgments.

The latest empirical studies also address gender biases in computers as algorithmic peers in screensaver tests. To that, (Storm et al., 2023) state: "Algorithmic screening devices will penalize women while comparing features or employment histories that are distinctly different from the so-called conventional, rigidly male career schemes." Even such measures in success found in the algorithm-based models hardly tend to accommodate those middle-of-the-road career paths defined by care-giving duties, which have a relative burden on women. Differences that do not stem from any positive gender-related attribute but seem functional take off as gendered patterns

of human experience. The way training and validation tasks are carried out on models greatly affects how extensive and long-lived such biases become. Findings from research shed some insights into how statistical modeling practices create gender-based outcomes in employment.

Racial bias in recruitment systems based on AI has also exhibited parallel trends in indirect discrimination emanating from model design and data choice. (Kumar, 2024) outline how technologies that use video analysis, language modeling, or speech recognition can provide unequal performance along racial lines. The use of facial analysis in computerized interviews, for instance, has shown lower accuracy for people with darker skin tones, and for women of color in particular. These differences are often traced back to underrepresentation in training data or to poorly generalized algorithmic models over diverse populations. The research identifies deeper institutional dynamics and technical assumptions behind racially differentiated outcomes. The findings provide a context in which differences along racial lines become evident through algorithmic operations.

Algorithmic recruitment exhibited intersectional prejudices that demarcated a new level wherein race intensified the effect of gender, making the structure all the more disadvantaged. Of note, while earlier studies (Amrouni et al., 2023) reported a wide variation in error rates across demographic subgroups, they indicated that errors in digital screening tools that employed facial recognition technologies most affected Black women, while their lowest or negligible positive outcome was with white men. This could imply that the systems embedded intersecting forms of demographic bias. According to the authors, this kind of difference can arguably be traced to a balance in the training data and further result from a lack of demographic representation while developing and assessing the systems. These form a historical and structural investment in data infrastructures that are central to analyzing these disparities. It allows technological examination into how algorithmic recruitment tools embed and replicate existing power disparities.

2. Studies on Candidate Perceptions of AI in the Recruitment Process

As regards the perception of job candidates about AI technologies used in recruitment procedures, it emerges as an area of concern in human-computer interaction and organizational behavior. (Gonzalez et al., 2022) evidence that job candidates seem to carry ambivalent feelings toward the AI-driven tools: automated resume screening and algorithm scoring of video interviews. Some candidates appreciate the speed and supposed objectivity of these tools, while others are concerned about fairness, data privacy, and appeal opportunities against the decisions made. These concerns are magnified when organizations offer no clarity on the AI's workings, thus alienating or making the candidates distrustful of the whole process. The authors state that transparency is an important factor determining whether candidates will perceive AI as an aid in

the recruitment workflow or as an affliction that dehumanizes their chances. In addition, an individual's interpretation of AI's legitimacy must be defined by personal past experiences, acceptance of technology, and contextual enactment of AI within recruitment.

The degree to which candidates consider AI systems fair also depends on how candidates think their characteristics and qualifications are considered. According to (Pergantis et al., 2025) many applicants would favor recruitment procedures over which humans judge, particularly for jobs requiring social interaction or creativity. AI carries the perception that it is a replacement for human decision-making instead of being a support mechanism, and it is seen as overly rigid or incapable of comprehending the finer nuances of a candidate's background. To form such perceptions, the output of the system is not the only factor; employers' communication around the role of AI also determines perceptions. Lack of explanation around the way AI comes to its conclusions is often enough for a candidate to consider the results arbitrary or impersonal. And thereby build up the degree of trust candidates have in the employer and their continued engagement in the recruitment process.

The literature delves deep into the candidates' emotional reactions during AI evaluations. (Li et al., 2025) show that anxiety, uncertainty, and frustration during AI assessments are very common, especially when the candidate is not provided any feedback or clarification. These emotional discomforts mainly arise from a lack of social cues, personal interaction, and real-time communication that are expected in a recruitment context. This impersonal experience leaves candidates feeling powerless and emotionally detached, particularly in high-stakes situations affecting outcomes that most certainly hold grave consequences. The study highlights that these emotional reactions are, therefore, contingent not only upon the smooth functioning of the AI systems but also on the expectations of the candidates that the recruitment experience would be fair and human-centered. Thus, emotional responses become an important avenue to analyze the ethics of AI use in hiring.

Cultural and contextualized factors contribute greatly to the ways analyzed by candidates when it comes to the fairness and objectivity attached to the AI tools of recruitment. Society norms on technology, data privacy, and institutional trust have deep and lasting effects on how algorithmic decision-making is received within societies (Starke et al., 2022). In this cross-national study, they found a great divergence in attitudes of candidates towards algorithmic decision making depending on established local norms about organizational accountability and civilian confidence in automated systems. In some countries, for instance, individuals conceive AI as more impartial and less biased than a human decision maker. Contrarily, historical backgrounds of evictions or visible opacity have heightened public skepticism about automation. Although it has emerged that

candidate perceptions are not static but alter along some variations in their individual opinions, how their organizational styles and employ their AI systems shape this perception. Therefore, all this implies that the candidate's perceptions are social and cultural instead of a uniform response assumed across various locations to the technology.

Table 1. Comparison of Previous Studies on Candidate Perceptions of AI in the Recruitment Process

Study	Focus of Analysis	Key Findings
(Gonzalez et al., 2022)	General attitudes toward AI-driven hiring tools	Candidates expressed both appreciation for efficiency and concern over fairness, privacy, and lack of transparency.
(Pergantis et al., 2025)	Perceived fairness and human discretion in AI-based selection	Preference for human involvement, especially in roles requiring nuanced judgment; AI seen as rigid and impersonal.
(Li et al., 2025)	Emotional reactions to AI assessments	Candidates often experienced anxiety and psychological detachment due to a lack of social cues and feedback.
(Starke et al., 2022)	Cross-cultural interpretation of AI fairness and transparency	Cultural norms influenced acceptance or skepticism of AI systems based on local trust and expectations of fairness.

III. RESEARCH METHOD

The research conducted in this study is based on a qualitative, exploratory document analysis to explore the ethical considerations of AI-driven recruitment platforms. Particularly about materials that cannot be quantified, are context dependent, and do not lend themselves to statistical generalizability, the qualitative paradigm is particularly appropriate for obtaining interpretive insights over statistical generalizability. With such an approach, one can grasp organizational documentation defining such ethical principles as fairness, transparency, and accountability. This study surveys specific sources-from white papers and policy documents to platform guidelines, elucidating the ethical commitments and perceived risks of AI recruitment systems. Simultaneously, it attempts to expose implicit logics and assumptions disparate to digital hiring practices, yet which may otherwise go unnoticed by quantitative inquiry. This interpretative analysis, then, acts as a starting point for exploring how various platforms use language and framing to signal their commitment to an ethical code and sway public and regulatory interpretation. In this respect, the research aims to highlight the recurrent themes and discursive strategies made regarding algorithmic bias, procedural transparency, and fairness as observed through the publicly available documents of different AI recruitment service providers.

The study draws upon various sources of data, including official ethical statements and corporate white papers from major recruitment platforms such as LinkedIn, HireVue, Pymetrics, and ModernHire, as well as publicly available documents such as Terms of Service, Privacy

Policies, and AI Ethics Statements. Apart from corporate literature, the study considers peer-reviewed academic literature and regulatory recommendations on artificial intelligence-based recruitment that provide not only important critical but also normative perspectives. These sources have been chosen as they serve to document technical and ethical concerns expressed by regulating bodies and platform developers. The fact that these sources are publicly accessible guarantees that their contents reflect official positions as communicated externally to interested parties, including job seekers, regulators, and researchers. This also serves to add a layer of context from regulatory and academic literature as to how industry practice would align or otherwise diverge from general ethical principles and academic critique, towards a fuller account of the discourse. All data were collected from the platform websites and publicly accessible digital archives, thereby enhancing the transparency and traceability of the research process with the possibility for future replication or exploration.

To provide a comprehensive overview of the primary documents analyzed, Table 2 below outlines the characteristics of the selected sources, including platform names, document types, year of publication, and the degree to which ethical considerations are explicitly addressed. This table serves as a reference point to illustrate the diversity of documentation available across digital hiring platforms and highlights variations in how each platform publicly communicates its ethical commitments. By comparing the presence or absence of explicit ethical content, the table allows readers to observe differences in how transparency and fairness are prioritized or framed across providers. This categorization also enables thematic comparison across platforms, particularly in terms of the clarity, scope, and depth of ethical articulation. Differences in documentation type, ranging from policy FAQs to technical audit reports, reveal not only content variation but also the intended audiences and communicative strategies employed. The information provided in the table forms a foundational dataset for the subsequent stages of thematic analysis, including text mining and manual coding of ethical themes, which aim to uncover the deeper patterns within ethical narratives and their operational implications.

Table 2. Characteristics of Analyzed Documents

Platform	Document Type	Year	Explicit Ethical Focus
LinkedIn	AI Ethics Statement & FAQ	2023	Yes (fairness & bias)
HireVue	Candidate Assessment Overview	2022	Partial
Pymetrics	AI Audit Report + Terms	2021	Yes (transparent auditing)
ModernHire	Whitepaper Predictive Hiring	2024	Not explicitly mentioned

There were three stages in the analysis process that sought to obtain ethical insights from the chosen documents. The first phase was dedicated to gathering documents from institutional websites and reputable public archives to confirm the content's validity and the trustworthiness of its sources. This step focused on capturing institutional 'whitelists' and self-regulatory

algorithms where recruitment was seen from an institutional viewpoint. In the second phase, text mining with AI was performed to recover phrases, analyze speech patterns, and classify topical themes using high-level natural language processing, including text-based GPT and Claude models. This analysis facilitated clearer identifications of surfacing concepts of fairness, accountability, auditability, and inclusivity that helped the research discover consistencies and discrepancies across platforms. In the last stage, thematic content analysis was performed through manual coding to extract core ethical themes such as implicit bias, black-box opacity, candidate control, and procedural justice. The cross and triangulated analysis from various document types within each ethical theme category ensured inter-interpretative validity and reinforced the trustworthiness of the ethical patterns that were uncovered.

This practical model enables an exploration from multiple angles as to how digital hiring platforms articulate and operationalize their ethical commitments. The method combines qualitative document examination with AI text mining and coding to facilitate both macro and micro-scale content analysis, addressing large-scale patterns as well as detailed contextual interpretations. These methods aid in comprehensive assessments of the ethics derived from corporate and legal documents and statements, which are bound to be either vaguely or clearly defined. The model is especially helpful in trying to determine the little-known ethical omissions and inclusions in stories regarding algorithm-based hiring. Also, examine the gap between ethical statements and working documents relative to various platforms. With this multi-layered analytical framework, the study aims to understand where algorithmic censuses fall concerning value-laden issues of balance, equity, fairness, power, responsibility, and dignity in the context of AI hiring systems.

IV. RESULT

A. Results

This study's analysis of platform documentation reveals a series of recurring ethical concerns in the implementation of AI-driven recruitment systems. Through qualitative coding and AI-assisted text analysis of official documents, including AI ethics statements, white papers, and privacy policies, from platforms such as LinkedIn, HireVue, Pymetrics, and ModernHire, several thematic categories emerged. These themes relate specifically to how these platforms articulate fairness, transparency, and accountability, and how such ethical concepts are operationalized, if they are implemented at all, within their hiring technologies. The data reveal that although platforms often employ similar terminology in their public-facing documents, the substance and clarity with which they address these ethical concerns vary considerably. In many cases, ethical principles are referenced without a clear articulation of how they are integrated into system design or decision-

making processes. This suggests that the ethical narratives presented in documentation may function more as symbolic affirmations than as concrete indicators of ethical practice in the development and deployment of AI recruitment tools.

The first central theme to emerge from the data is the varying emphasis on ethical principles across platforms, as indicated by the frequency of key ethical terms found in the analyzed documents. Ethical vocabulary such as fairness, bias, transparency, explainability, and accountability appeared with different degrees of frequency across platforms, reflecting inconsistencies in how ethical considerations are prioritized. This pattern is demonstrated in Figure 1, a heatmap that visualizes the relative frequency of these ethical terms across each platform. The heatmap shows that LinkedIn and HireVue use ethical terminology more frequently than Pymetrics and ModernHire, particularly concerning fairness and bias. Such differences in the presence of ethical language may reflect divergent institutional priorities, strategic branding approaches, or varying levels of investment in ethical AI development. The findings underscore the need to critically examine not only the presence of ethical language but also the context in which such language is employed, as it may serve both normative and reputational functions.

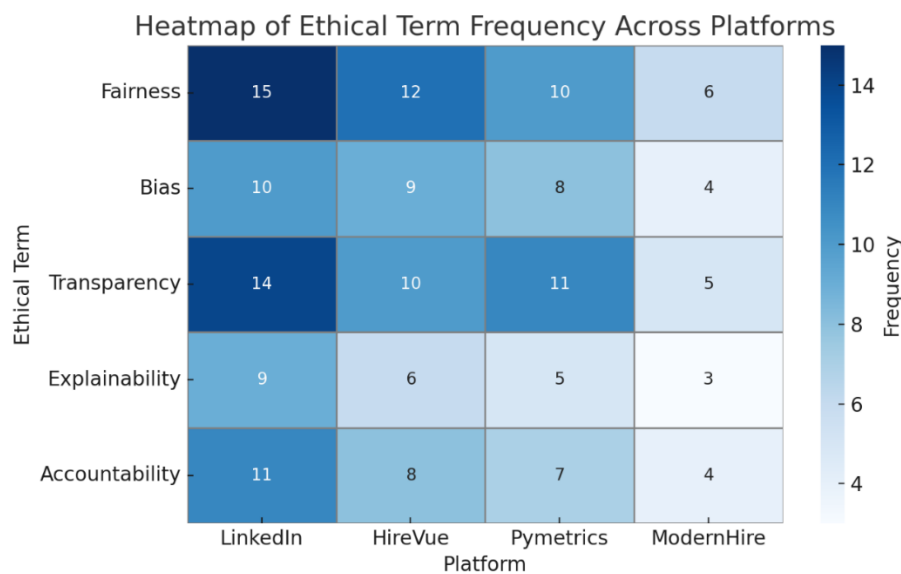


Figure 1. Heatmap of Ethical Term Frequency Across Platforms

The second theme focuses on the systematic approach to clarity and comprehensiveness concerning specific statements of transparency, especially about the operations of AI-enabled hiring systems. While platforms often mention the use of AI for efficiency and predictive prowess, they tend to skip providing ample details on the processes involved in candidate data extraction, systematized decision making, or how the constituent components of the system interact. This gap in information poses a problem for job seekers who do not have sufficient knowledge

regarding the processes performed on their profiles and how they are evaluated or ranked. To demonstrate these gaps, Figure 2 provides a comparison of the transparency score across platforms, which considers the presence and adequacy of transparency-related information in publicly available documents. It can be seen that in terms of transparency, LinkedIn emerged on top, followed by HireVue, Pymetrics, and ModernHire in that order. The pattern of these scores shows that although a degree of value voting is undertaken, as expressed by the overwhelming focus on sophistication, there is an evident gap in how platforms implement this claim in practice, especially in their public statements and disclosures regarding algorithmic logic.

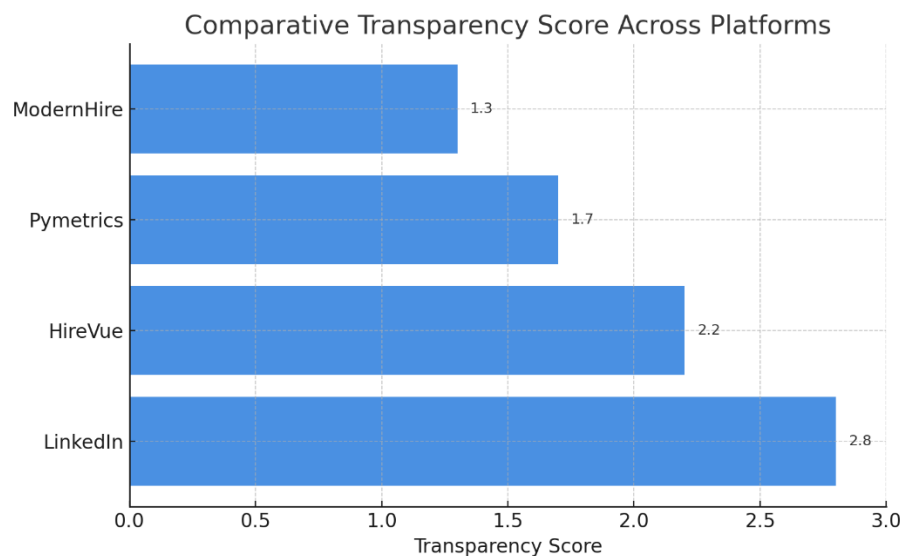


Figure 2. Comparative Transparency Score Across Platforms

Outside of the terminology and transparency issues, the examination revealed some previously noted ethical gaps stemming from the algorithms and governance of AI-based hiring systems. These include underlying biases within automated systems, absent candidate rebuttal processes, and gaps in automation explainability. As this suggests, the issues are not singular; rather, they are part of a cascading chain of ethical challenges that tend to impede candidates within the recruiting process. These three areas of concern emerged more strongly at the thematic coding stage and were triangulated through various documents, including technical documents, ethics policies, and policy documents. Take, for example, bias. It may stem from an unfortunate reliance on the training data, which contains some form of social or historical discrimination and therefore reinforces such inequitable practices. In the case of opacity, the justification involves limited disclosure of algorithmic frameworks, and the sheer sophistication of the AI systems combats scrutiny or authoritative challenge from external parties. For this reason, concern is acute for the lack of candidate mechanisms to appeal or obtain an audit of the decision made about them, and permitting candidates the discretionary power becomes pivotal in establishing trust within the

system. These findings are summarized in Table 3, which outlines the dominant ethical themes, their sub-descriptions, and the platforms where they are most frequently encountered.

Table 3. Key Ethical Themes in AI Recruitment Systems

Main Theme	Subtheme / Description	Dominant Platforms
Algorithmic Bias	Implicit preferences related to gender or race	HireVue, LinkedIn
Opaque Process	Lack of clarity regarding how candidates are evaluated	All platforms
Lack of Recourse	Absence of appeals or clarification mechanisms	LinkedIn, ModernHire

Collectively, these findings illustrate the ethical scope of AI recruiting systems. They show how ethical considerations, while present, are unevenly implemented across platforms and how critical issues of bias, transparency, and candidate empowerment still linger in the crafting and discourse of algorithmic hiring technologies. The analysis indicated that even in cases where ethical considerations are at least superficially acknowledged, the enforcement of such values tends to be incoherent, shallow, and performative. Digital recruitment platforms offer a varied landscape in their approaches toward ethical governance, some providing well-formed frameworks and others simply limiting themselves to loosely defined aspirational rhetoric rarely applied. The extent of meaningful oversight through elements such as control mechanisms, accountability protocols, and opportunities for user participation seems contingent upon internal considerations of the organization and on external factors, such as public scrutiny and regulatory engagement. What comes through is a rather fragmented and uneven landscape of ethics, where reactive and sometimes stifling responses to institutional pressure are tempered by innovative efforts to shape the standards in development. Rather than forming some form of cohesive governance, many platforms appear to be addressing their ethical responsibilities in an ad hoc manner that serves their reputations just as much as it protects their users.

V. DISCUSSION

Results from the ongoing research indicate that there exists an almost forever gap between professed ethical principles regarding digital recruitment platforms and actual practices. Though some of these big platforms- HireVue, LinkedIn- profess some values, such as fairness, transparency, or explainability, there is always an element of doubt about whether or not these values remain committed in spirit. This is in line with (Tavares & Ferrara, 2023) and (Ricart et al., 2022), where it has been considered that ethical declarations within algorithmic hiring are more often than not a smokescreen for painting over the existing realities rather than being operational commitments. Such contrasts would apply per how this ethical lingo has been applied unevenly across platforms; oftentimes, references to fairness or transparency fail to translate in a practical sense into governance structures. Current observations corroborate the assertion made by (Thomaidou & Limniotis, 2025), which states that there can be no meaningful guarantees for

transparency in AI systems if there are no genuine and relatable avenues for seeking explanation by end-users. Similar findings reaffirm (Ortega-Bolaños et al., 2024) in arguing that ethical values ought to be ingrained slithery within and not be superficially appended thereafter in the design process. Quite truly then, with the process of decision-making forever being obscure and with scanty systems of accountability in place, it can be stated that performance and efficiency trump integrity in the context of the development of AI applications for recruitment.

Apart from validating previous scholarships, the study makes further contributions to the current understanding of how ethical risks like algorithmic bias and lack of candidate recourse become structurally rooted in the platforms. According to (Amrouni et al., 2023), algorithmic bias is often generated as a result of the bias in training data and weak auditing of system architecture. These are not just flaws but systemic oversights, especially when it comes to ways that fairness constraints are unevenly applied. The disclaimer of how a decision is taken will bring in the accentuation of worries as suggested by (Hassija et al., 2024), about the so-called "black-box" paradigm of AI systems. Most candidates have no clue about which data points could decide the outcome, develop biases, develop suspicion, as well as a lack of accountability. In addition, it lacks mention of independent audits or procedural oversight, which has also been raised under (Felländer et al., 2022). It underlines, therefore, a larger failure in setting up ethical safeguards. Findings show that most of the platforms do not have clear appeal procedures either, thus further limiting the agency of candidates and reinforcing even more the power imbalance between users and systems. Together, these trends suggest that ethical AI development in recruitment will remain inconsistent, unapparent, and even unjust unless stronger and more effective regulatory frameworks and internal accountability measures are put in place.

VI. CONCLUSION AND RECOMMENDATION

The study AI Ethics concluded that problems with bias and other ethical aspects, like transparency, still plague the algorithmic hiring tools. The gap between fairness and accountability and what is being done in practice is a common feature across digital hiring systems. While they claim to base systems on certain ethical principles, it appears that such principles are only paid lip service to. One of the most serious issues is the lack of appeal mechanisms in place to challenge the decisions that algorithms make or the decisions in the process. These flaws pose fundamental problems on whether candidates as not equal and autonomous participants in the employment process. The algorithms themselves frequently reflect and reproduce basic social inequalities alongside historical discrimination, which renders these technologies far from neutral, employing preexisting social hierarchies. Without robust auditing measures, these systems have few checks and balances, making them prone to misuse while

escaping accountability easily. Federated learning implies stronger data privacy, but does not eradicate the requirement for centralization. These findings suggest that there is a need to imbue ethics in the design and governance process during the AI life cycle of recruitment.

A response to these concerns, HR technology firms need to take responsible and transparent stances on AI in hiring. This involves explainable building systems, such that applicants or recruiters know how decisions are made. Make sure also to have clear review and appeal procedures that allow candidates to challenge or appeal results. The time is ripe for the development of international ethical standards for the use of AI in recruitment, and indeed, ethical considerations should play a central role in the design and monitoring of AI for selection to prevent discrimination and ensure fairness across social and geographic boundaries. Next steps for research include more grounded, empirical forms of research, for example, in-depth interviews with recruiters and applicants, that can capture how AI systems are experienced in practice in hiring (city). Those studies would help shed light on the emotional and social aspects of algorithmic hiring, aspects that are easily dismissed if all one looks at are technical tools. Furthermore, the ethical imperatives of fairness and transparency also must be viewed in the context of cultural differences since the expectations and values vis-à-vis technology could be different across regions. This is the only way that we will develop a kind of recruitment technology that serves human needs - and ultimately ethical objectives - at every level.

REFERENCES

- Almeida, F., Junça Silva, A., Lopes, S. L., & Braz, I. (2025). Understanding Recruiters' Acceptance of Artificial Intelligence: Insights from the Technology Acceptance Model. *Applied Sciences*, 15(2), 746. <https://doi.org/10.3390/app15020746>
- Amrouni, N., Benzaoui, A., & Zeroual, A. (2023). Palmprint Recognition: Extensive Exploration of Databases, Methodologies, Comparative Assessment, and Future Directions. *Applied Sciences*, 14(1), 153. <https://doi.org/10.3390/app14010153>
- Balasubramaniam, N., Kauppinen, M., Rannisto, A., Hiekkänen, K., & Kujala, S. (2023). Transparency and Explainability of AI Systems: From Ethical Guidelines to Requirements. *Information and Software Technology*, 159, 107197. <https://doi.org/10.1016/j.infsof.2023.107197>
- Chen, P., Wu, L., & Wang, L. (2023). AI Fairness in Data Management and Analytics: A Review on Challenges, Methodologies and Applications. *Applied Sciences*, 13(18), 10258. <https://doi.org/10.3390/app131810258>
- Eke, D., & Stahl, B. (2024). Ethics in the Governance of Data and Digital Technology: An Analysis of European Data Regulations and Policies. *Digital Society*, 3(1), 1–23. <https://doi.org/10.1007/s44206-024-00101-6>
- Felländer, A., Rebane, J., Larsson, S., Wiggberg, M., & Heintz, F. (2022). Achieving a Data-Driven Risk Assessment Methodology for Ethical AI. *Digital Society*, 1(2), 13.

<https://doi.org/10.1007/s44206-022-00016-0>

- Gonzalez, M. F., Liu, W., Shirase, L., Tomczak, D. L., Lobbe, C. E., Justenhoven, R., & Martin, N. R. (2022). Allying With AI? Reactions Toward Human-Based, AI/ML-Based, and Augmented Hiring Processes. *Computers in Human Behavior*, 130, 107179. <https://doi.org/10.1016/j.chb.2022.107179>
- Hassija, V., Chamola, V., Mahapatra, A., Singal, A., Goel, D., Huang, K., Scardapane, S., Spinelli, I., Mahmud, M., & Hussain, A. (2024). Interpreting Black-Box Models: A Review on Explainable Artificial Intelligence. *Cognitive Computation*, 16(1), 45–74. <https://doi.org/10.1007/s12559-023-10179-8>
- Kumar, P. (2024). Large Language Models (LLMs): Survey, Technical Frameworks, and Future Challenges. In *Artificial Intelligence Review* (Vol. 57, Issue 9). Springer Netherlands. <https://doi.org/10.1007/s10462-024-10888-y>
- Lamers, L., Meijerink, J., Jansen, G., & Boon, M. (2022). A Capability Approach to Worker Dignity Under Algorithmic Management. *Ethics and Information Technology*, 24(1), 1–15. <https://doi.org/10.1007/s10676-022-09637-y>
- Li, X., Huang, M., Liu, J., Fan, Y., & Cui, M. (2025). The Impact of AI Negative Feedback vs. Leader Negative Feedback on Employee Withdrawal Behavior: A Dual-Path Study of Emotion and Cognition. *Behavioral Sciences*, 15(2), 152. <https://doi.org/10.3390/bs15020152>
- Munifah, M., Wibawa, E. S., & Purwantini, K. (2024). Ethical Challenges in AI-Driven Decision-Making: Addressing Bias and Accountability in Business Applications. *Journal of Management and Informatics*, 3(1), 105–121. <https://doi.org/10.51903/jmi.v3i1.48>
- Nannini, L., Marchiori Manerba, M., & Beretta, I. (2024). Mapping the Landscape of Ethical Considerations in Explainable AI Research. *Ethics and Information Technology*, 26(3), 44. <https://doi.org/10.1007/s10676-024-09773-7>
- Ortega-Bolaños, R., Bernal-Salcedo, J., Germán Ortiz, M., Galeano Sarmiento, J., Ruz, G. A., & Tabares-Soto, R. (2024). Applying the ethics of AI: a systematic review of tools for developing and assessing AI-based systems. *Artificial Intelligence Review*, 57(5), 1–30. <https://doi.org/10.1007/s10462-024-10740-3>
- Pagano, T. P., Loureiro, R. B., Lisboa, F. V. N., Peixoto, R. M., Guimarães, G. A. S., Cruz, G. O. R., Araujo, M. M., Santos, L. L., Cruz, M. A. S., Oliveira, E. L. S., Winkler, I., & Nascimento, E. G. S. (2023). Bias and Unfairness in Machine Learning Models: A Systematic Review on Datasets, Tools, Fairness Metrics, and Identification and Mitigation Methods. *Big Data and Cognitive Computing*, 7(1), 1–31. <https://doi.org/10.3390/bdcc7010015>
- Pergantis, P., Bamicha, V., Skianis, C., & Drigas, A. (2025). AI Chatbots and Cognitive Control: Enhancing Executive Functions Through Chatbot Interactions: A Systematic Review. *Brain Sciences*, 15(1), 47. <https://doi.org/10.3390/brainsci15010047>
- Ricart, C. P., Rodrigues, J. M. F., Cardoso, P. J. S., Varona, D., & Luis Suárez, J. (2022). Discrimination, Bias, Fairness, and Trustworthy AI. *Applied Sciences*, 12(12), 5826. <https://doi.org/10.3390/app12125826>
- Rigotti, C., & Fosch-Villaronga, E. (2024). Fairness, AI & recruitment. *Computer Law and Security Review*, 53, 105966. <https://doi.org/10.1016/j.clsr.2024.105966>
- Starke, C., Baleis, J., Keller, B., & Marcinkowski, F. (2022). Fairness Perceptions of Algorithmic

Decision-Making: A Systematic Review of the Empirical Literature. *Big Data and Society*, 9(2), 20539517221115188. <https://doi.org/10.1177/20539517221115189>

Storm, K. I. L., Reiss, L. K., Guenther, E. A., Clar-Novak, M., & Muhr, S. L. (2023). Unconscious Bias in the HRM Literature: Towards a Critical-Reflexive Approach. *Human Resource Management Review*, 33(3), 100969. <https://doi.org/10.1016/j.hrmr.2023.100969>

Sumarlin, T., & Kusumajaya, R. A. (2024). AI Challenges and Strategies for Business Process Optimization in Industry 4.0: Systematic Literature Review. *Journal of Management and Informatics*, 3(2), 195–211. <https://doi.org/10.51903/jmi.v3i2.25>

Tavares, S., & Ferrara, E. (2023). Fairness and Bias in Artificial Intelligence: A Brief Survey of Sources, Impacts, and Mitigation Strategies. *Sci*, 6(1), 3. <https://doi.org/10.3390/sci6010003>

Thomaidou, A., & Limniotis, K. (2025). Navigating Through Human Rights in AI: Exploring the Interplay Between GDPR and Fundamental Rights Impact Assessment. *Journal of Cybersecurity and Privacy*, 5(1), 7. <https://doi.org/10.3390/jcp5010007>

Tursunalieva, A., Alexander, D. L. J., Dunne, R., Li, J., Riera, L., & Zhao, Y. (2024). Making Sense of Machine Learning: A Review of Interpretation Techniques and Their Applications. *Applied Sciences (Switzerland)*, 14(2), 496. <https://doi.org/10.3390/app14020496>

Wahyuning, S., & Sudibyo, S. K. (2024). Leveraging Machine Learning for Talent Acquisition: Predicting High-Performance Candidates in Human Resource Management. *Journal of Management and Informatics*, 3(1), 87–104. <https://doi.org/10.51903/jmi.v3i1.44>