



LUND UNIVERSITY

School of Economics and Management

Department of Informatics

Efficient, but Ethical?

A qualitative study on AI-supported recruitment

Bachelor thesis, 15 hp, kurs SYSK16 in Informatics

Authors: Hannah Grundström
Elsa Högberg
Ines Kahrmanovic

Supervisor: Osama Mansour

Examiners: Nicklas Holmberg
Umberto Fiaccadori

Efficient, but Ethical? A Qualitative Study on AI-supported Recruitment

AUTHORS: Hannah Grundström, Elsa Högberg, Ines Kahrmanovic

PUBLISHER: Department of Informatics, School of Economics and Management, Lund University

FORMAL EXAMINATOR: Osama Mansour, Docent

PRESENTED: May, 2026

DOCUMENT TYPE: Bachelor thesis

NUMBER OF PAGES: 49

KEYWORDS: AI ethics, algorithmic decision making, human judgement, recruitment, HR-professionals

SUMMARY (MAX. 200 WORDS):

The use of AI-supported recruitment systems is increasing rapidly, raising important ethical questions. The study explores how HR professionals perceive and manage ethical challenges related to AI-supported recruitment, as well as how these challenges are managed in practice. The study is based on a qualitative method where semi-structured interviews were conducted with six respondents working with AI-supported recruitment in Sweden and the United States. The analysis is grounded in a theoretical framework and previous research concerning algorithmic bias, fairness, transparency, human oversight, and AI regulation. The findings indicate that AI is often perceived as a valuable support tool that may contribute to more structured and consistent recruitment processes while reducing certain forms of subjective human bias. At the same time, respondents express concerns regarding algorithmic bias, lack of transparency, overreliance on automated systems and uncertainty regarding the interpretation of legal and regulatory requirements. The study further shows that human involvement, education, and critical evaluation of AI-generated outputs are perceived as central strategies for mitigating ethical risks. Overall, the findings suggest that ethical AI use in recruitment is shaped through the interaction between technological systems, human judgment, organizational practices, and the interpretation of legal and regulatory frameworks.

Table of contents

1 Introduction.....	1
1.1 Background.....	1
1.2 Problem statement.....	2
1.3 Purpose.....	3
1.4 Research Question.....	3
1.5 Delimitations.....	3
2 Literature review.....	5
2.1 AI in recruitment.....	5
2.1.1 Algorithmic decision-making.....	5
2.1.2 Algorithmic decision-making in recruitment and its benefits.....	6
2.2 Ethical concerns in AI recruitment.....	7
2.2.1 Bias and discrimination.....	7
2.2.2 Algorithmic Fairness.....	8
2.2.3 Transparency.....	10
2.2.4 Accountability.....	10
2.3 Legal and regulatory considerations in recruitment.....	11
2.3.1 Transparency and regulatory challenges.....	11
2.3.2 Accountability and human oversight.....	12
2.4 Mitigating Ethical Risks in AI-Driven Recruitment.....	13
2.4.1 Transparency and Explainability Measures.....	13
2.4.2 Bias Detection and Fairness-aware Models.....	13
2.4.3 Human Oversight in Decision-making.....	14
2.4.4 Governance, Policies, and Regulatory Compliance.....	14
3 Theory.....	15
3.1 Ethical Framework for AI-enabled recruiting.....	15
3.2 Application of the Ethical Framework.....	15
3.2.1 Refinement of Framework Themes.....	16
4 Methodology.....	18
4.1 Qualitative research.....	18
4.2 Semi-structured interviews.....	18
4.2.1 Sampling approach.....	19
4.2.2 Design of interview guide.....	20
4.2.3 Interview preparation.....	20
4.2.4 Data collection procedure.....	21
4.3 Analysis.....	21

4.3.1 Transcription.....	21
4.3.2 Thematic analysis.....	23
4.4 Reliability & validity.....	24
4.4.1 Ethical considerations.....	25
5 Empirical data.....	26
5.1 Ethical opportunities.....	26
5.1.1 Efficiency and process support.....	26
5.1.2 Reduction of human bias.....	27
5.2 Ethical risks.....	27
5.2.1 Introduction to algorithmic bias.....	27
5.2.2 Lack of transparency.....	28
5.2.3 Potential loss of human oversight and accountability.....	28
5.2.4 Legal and ethical risks related to data usage.....	29
5.2.5 Assessment limitations and misinterpretation.....	29
5.3 Ethical ambiguities.....	30
5.3.1 Effect on workforce diversity.....	30
5.3.2 Impact on assessment validity and accuracy.....	30
5.3.3 Perceived fairness.....	31
5.3.4 The interpretation and application of AI regulation.....	31
5.4 Mitigation strategies.....	32
6 Discussion.....	34
6.1 Ethical opportunities.....	34
6.1.1 Efficiency and process support.....	34
6.1.2 Reduction of human bias.....	35
6.2 Ethical risks.....	36
6.2.1 Introduction to algorithmic bias.....	36
6.2.2 Lack of transparency.....	36
6.2.3 Potential loss of human oversight and accountability.....	37
6.2.4 Legal and ethical risks related to data usage.....	38
6.2.5 Assessment limitations and misinterpretation.....	38
6.3 Ethical ambiguities.....	39
6.3.1 Effect on workforce diversity.....	39
6.3.2 Perceived fairness.....	40
6.3.3 The interpretation and application of AI regulation.....	41
6.4 Approaches to mitigate ethical risks.....	42
7 Conclusion.....	43
7.1 Suggestions for Further Research.....	43



LUND UNIVERSITY

School of Economics and Management

Figures

Figure 1: Ethical Framework of AI-Enabled Recruiting (Hunkenschroer & Luetge, 2022).... 15

Tables

Table 1. Ethical considerations.....	16
Table 2. Summary of Respondents.....	19
Table 3. Coding scheme for data analysis.....	23
Table 4. Extract from Coding Scheme.....	24

1 Introduction

1.1 Background

Artificial intelligence (AI) is becoming increasingly integrated into the recruitment processes, as organizations want to automate tasks that previously were done manually. This development has become more prominent since organizations receive larger volumes of applications, where they are expected to streamline their hiring process and make recruitment decisions more efficiently (Black & Van Esch, 2020; Varshney, 2025).

In recruitment, AI tools can be used throughout the whole hiring process, automating various parts of recruitment (Yanamala, 2021). Such systems may support outreach by identifying potential candidates and targeting job advertisements. They may also support screening by parsing and ranking resumes, assessment through chatbot- or video-based tools, and facilitation by automating administrative tasks such as scheduling and communication (Black & Van Esch, 2020; Hunkenschroer & Luetge, 2022). It does not work as a single tool, but rather as a set of technologies integrated in different parts of the process.

Organizations are motivated to use AI in recruitment due to its benefits. The systems can process and analyze large amounts of data, identify patterns, and generate outputs more quickly and consistently than humans alone can do. In this sense, AI is often presented as a way to improve efficiency, support decision-making, and create more standardized recruitment practices. It may also reduce subjective human judgment, making evaluations more objective (Köchling & Wehner, 2020; Logg et al., 2019).

Nevertheless, the use of AI in recruitment has become progressively more debated. Research finds that these systems do not simply remove bias, but may actually reproduce existing bias depending on how the systems are designed or trained (Köchling & Wehner, 2020; Möller & Petrovic, 2025). Additionally, the algorithmic systems are often opaque, and recruiters may struggle to understand or explain the produced outputs (Möller & Petrovic, 2025; Zhang et al., 2025). This raises concerns regarding fairness, transparency, accountability, and the role of human judgment in AI-supported recruitment processes.

Overall, the growing integration of AI in recruitment reflects a shift toward more data-driven and automated hiring processes, driven by the need for efficiency and scalability in handling large applicant volumes. As mentioned, there are many advantages such as speed and consistency, but it also comes with complexities like bias, transparency, and accountability. As AI becomes embedded across multiple stages of the recruitment process, it should be viewed

as a set of interconnected technologies that interact with human judgment and organizational practices, rather than as a single neutral tool. Therefore, understanding how these systems function in practice has become increasingly important.

1.2 Problem statement

Despite these developments, research also finds several ethical challenges. Studies have found that algorithmic systems may reproduce existing biases present in historical data, leading to discriminatory outcomes (Köchling & Wehner, 2020). Some studies have found instances of racial bias (Fabris et al., 2026; Pager et al., 2009), while others have identified instances of gender bias (Carlsson et al., 2021; Sony et al., 2025). At the same time, the complexity and opacity of algorithmic systems, often described as the “black box” problem, make it difficult for recruiters to understand, explain, or justify the algorithmic decisions. This influences how these systems are used in practice. It raises concerns about automation bias and diffused accountability, as human decision-makers may over-rely on algorithmic recommendations (Karami & Ghasemi, 2026; Leicht-Deobald et al., 2019).

In addition, research suggests that fairness in AI systems is not only a technical issue, but also involves trade-offs between competing objectives such as accuracy, non-discrimination, and organizational goals (Teo, 2025). While efforts have been made to improve algorithmic fairness, empirical findings indicate that fairer algorithmic designs do not necessarily lead to fair outcomes in practice. Instead, recruiters interpret and act upon algorithmic outputs, which are also influenced by contextual factors, such as interface design and time constraints (Fabris et al., 2026)

Prior studies have examined issues such as algorithmic bias, transparency, and fairness, (Kumaresh et al., 2025; Yanamala, 2023) and have begun to explore the interaction between AI systems and human decision-makers (Laukkarinen, 2025; Wang et al., 2025). However, the research in this area is relatively limited and fragmented. There is still a limited understanding of how ethical challenges emerge in practice when recruiters interact with AI systems in recruitment, especially as organizations continue to adopt and use AI systems. This creates a tension between the perceived benefits of AI and its potential negative consequences.

This gap is significant because recruitment decisions have consequences for individuals' future job opportunities (Hunkenschroer & Luetge, 2022). Decisions made at the initial stages of the recruitment process determine who is considered further and who is excluded, meaning that potential biases or ethical issues may have disproportionate effects on the whole process (Fabris et al., 2026). Furthermore, recent regulatory developments (e.g., the EU AI Act) and existing frameworks (e.g., the GDPR) also underscore the need for organizations to ensure that AI-supported recruitment processes are ethically and legally compliant in practice (Yanamala, 2023).

Therefore, this research investigates what ethical challenges emerge in AI-supported recruitment, particularly in the interaction between algorithmic systems and human decision-makers. This can contribute to a better understanding of how ethical and regulatory requirements are managed in practice, and how AI can be used more responsibly in recruitment.

1.3 Purpose

This study aims to explore and develop an understanding of the ethical implications of using AI in recruitment. Using a qualitative approach, the study examines how these challenges are experienced and handled by HR professionals.

The study focuses on how recruiters could balance algorithmic recommendations in situations where there are known risks, and how these are weighed against the perceived benefits of using AI. It also aims to provide a better understanding of how such challenges can be managed, and under what conditions AI can be used more responsibly.

The results are expected to contribute to insights relevant for HR professionals and organizations working with AI in recruitment. It will also benefit candidates by deepening their understanding of how AI is used in recruitment.

1.4 Research Question

Based on our problem statement, the following research question was formulated:

What ethical considerations emerge when using AI in recruitment, and how is it addressed by HR professionals?

1.5 Delimitations

This study is limited to the perspectives of HR professionals and individuals working with AI-supported recruitment systems. Consequently, candidates' own experiences and perceptions of AI-assisted recruitment processes are not examined. Furthermore, the study does not include the perspectives of AI developers, as its purpose is not to technically evaluate or compare specific AI systems or algorithms.

The research primarily focuses on recruitment contexts in which AI systems are used to assess relatively standardized application materials, such as resumes and written candidate information. Recruitment processes that rely heavily on creative portfolios or highly specialized practical assessments were therefore not examined in depth.

Geographically, the study includes respondents from Sweden and the United States. Swedish respondents operate within frameworks such as the GDPR and the EU AI Act, while respondents in the United States face a different regulatory environment. However, differences between these regulatory systems are not evaluated in detail within the scope of this study. In addition, since the study is based on a qualitative research approach and includes only six respondents, the findings are not intended to be statistically generalizable. The relatively small sample size also limits the ability to draw broader conclusions beyond the specific contexts represented in the study.

2 Literature review

2.1 AI in recruitment

2.1.1 *Algorithmic decision-making*

Algorithmic decision-making (ADM) refers to the use of algorithms to support or replace human decisions in organizations (Wang, 2025). An algorithm is a step-by-step procedure that takes inputs (e.g., data) and produces outputs (e.g., predictions or recommendations). AI is based on algorithms, and the systems are designed to perform tasks that would traditionally require human intelligence (Keegan & Meijerink, 2025; Krauss, 2024). This allows organizations to structure and standardize decision-making processes across different functions.

ADM can take different forms depending on its purpose. Some algorithms are descriptive (i.e., they record and organize past information), while others are predictive (i.e., they estimate what is likely to happen in the future). More advanced systems are prescriptive, meaning they recommend or automatically choose actions based on available data (Wang, 2025). While the different forms vary in complexity, in most cases they do not fully replace human decision-making (Keegan & Meijerink, 2025). This suggests that ADM should not be understood as purely automated decision-making. Decisions should be made by a combination of algorithms and human judgment.

Organizations have started to use ADM because it can significantly improve decision-making processes. AI can process and analyze large amounts of data, identify patterns, and produce consistent outputs. Algorithmic judgments appear to outperform human judgements in terms of accuracy (Logg et al., 2019; Yanamala, 2021). As a result, organizations use ADM in areas such as performance evaluation, resource allocation, and planning (Wang, 2025). Together, these advantages make ADM a powerful tool that enhances efficiency and consistency.

Despite the advantages, a key challenge is how people respond to the algorithmic systems. Research shows that human reactions to ADM may differ, characterized by a tension between algorithmic aversion and algorithmic appreciation (Dietvorst, 2016; Logg et al., 2019). On the one hand, individuals may prefer their own judgment over algorithmic advice. This is called algorithmic aversion, where people trust algorithms less because they are less forgiving of algorithmic mistakes than human mistakes, even when algorithms are more accurate (Dietvorst, 2016). On the other hand, studies have found evidence of algorithmic appreciation, where individuals place greater trust in and rely more heavily on algorithmic advice than human advice, even when they do not fully understand how the algorithm works (Logg et al.,

2019). Nevertheless, research also suggests that human-made decisions are often perceived as more trustworthy than AI-assisted or automated decision-making processes (Orbán & Stefkovics, 2025).

These mixed findings show that using ADM is not only about having more accurate tools, but also about how people understand, interpret, and use them in processes. Factors such as transparency, user control, and the type of task may influence whether people accept or reject algorithmic advice (Dietvorst, 2016; Logg et al., 2019). Therefore, ADM should not be seen only as a technical system, but also as something shaped by how people interact with it in organizations.

2.1.2 Algorithmic decision-making in recruitment and its benefits

Algorithmic decision-making is increasingly used in recruitment, where artificial intelligence is employed to support various stages of the process. The use of AI in recruitment offers several advantages, including improving efficiency and streamlining the hiring process, which can lead to faster decision-making (Varshney, 2025). Specifically, AI can process large volumes of information at a speed that exceeds human capability, allowing organizations to manage a greater number of applications more effectively. This may help organizations target suitable candidates while also attracting a more diverse pool of applicants (Black & Van Esch, 2020).

Two different types of AI can be used in Human Resource Management (HRM): discriminative AI and generative AI. Discriminative AI is primarily used to analyze and predict information (e.g., resume screening or candidate ranking systems). Generative AI supports the creation of content (e.g., generating job descriptions or assisting communication processes). These systems are integrated into recruitment practices, where AI and human decision-makers work together, described as a “*conjoined agency*” (Dutta & Naveen, 2025). AI is applied across several stages of recruitment, including outreach, screening, assessment, and facilitation (Hunkenschroer & Luetge, 2022).

In the *outreach stage*, organizations aim to identify and attract potential candidates. Although often overlooked, this phase has become increasingly important due to the rise of social media platforms, such as LinkedIn (Varshney, 2025). AI can help organizations identify and target suitable candidates more efficiently by screening online profiles and distributing job advertisements to relevant audiences (Alexander III et al., 2024). AI can also support communication, which improves candidate engagement and response rates, helping organizations reduce time-to-hire (Muhammad, 2025; Varshney, 2025).

In the *screening stage*, AI systems are used to filter, evaluate, and rank candidates based on predefined criteria. Tools such as automated resume screening and natural language processing improve the match between job requirements and applicant profiles (Hawrysz, 2025). AI can help make hiring more objective by using data to assess candidates. It can also

automate routine screening tasks, reduce the risk of human mistakes, and improve the accuracy of finding qualified applicants (Köchling & Wehner, 2020).

In the *assessment stage*, AI can be applied in both video- and chatbot-based interviews, as well as skills tests. These systems can analyze additional data points (e.g., tone of voice and facial expressions) to provide insights into candidates' skills and characteristics (Hunkenschroer & Luetge, 2022). Interview chatbots have been shown to screen large numbers of applicants efficiently, making the assessment process more scalable and standardized. However, these technologies are still relatively underdeveloped and less commonly used (Koivunen et al., 2022).

In the *facilitation stage*, AI supports the overall recruitment process by automating administrative tasks and improving communication. For example, AI can extract information from resumes, schedule interviews, and provide real-time updates to candidates, which increase efficiency and consistency (Hunkenschroer & Luetge, 2022).

These stages show that AI is integrated throughout the whole process and shapes how candidates are evaluated. Recruitment decisions are therefore made by the interaction between algorithms and human judgement. While this improves efficiency and consistency and it may also reduce the risk of human bias, it does raise questions about how these systems influence decision-making in practice, particularly from an ethical perspective. This is discussed in the following section.

2.2 Ethical concerns in AI recruitment

2.2.1 Bias and discrimination

AI systems are often introduced with the expectation that they can reduce human bias and promote more objective and consistent decision-making. However, research indicates that these systems may also introduce new forms of bias and discrimination. When humans rely on algorithmic decision-making, it poses a risk for discriminatory or unfair outcomes, since algorithms may be trained on inaccurate or biased data (Köchling & Wehner, 2020). Several studies find different kinds of biases, such as racial bias, gender bias and heterosexist bias.

Racial bias is defined as “*unfair judgment of a person based on their race*” (Cambridge Dictionary, n.d.a). In a study on contemporary discrimination, applicants applied to hundreds of jobs with resumes that were equal. Even though the candidates were equally qualified, the white candidates were twice as likely as the black candidates to receive a job offer. Black candidates with no criminal history did not even outperform white applicants who had recently been released from prison (Pager et al., 2009). The findings in Fabris et al.'s (2026) study, which took place in Western Europe, show that visible cues (e.g., foreign-sounding

names or education and work experience from other countries) hurt non-European candidates. When profiles signaled recent migration, their chance of being shortlisted fell by about 29% even though their qualifications were the same.

Another type of bias is gender bias, which is defined as “*the unfair difference in the way women and men are treated*” (Cambridge Dictionary, n.d.b). It remains a prevalent issue in recruitment, stemming from a long history of unequal treatment within the labor market. Such disparities are frequently reinforced by implicit stereotypes and assumptions about gender roles, which can influence recruiters’ judgments in subtle but systematic ways.

Similarly, a study done in the Nordic shows that women are largely underrepresented in professor positions. There is a large gap between the genders in high-level roles, reflecting broader structural inequalities within academia. Despite increasing numbers of women entering higher education and completing doctoral degrees, this representation does not translate proportionally into senior academic positions (Carlsson et al., 2021). To further illustrate these risks, Sony et al. (2025) refer to the case of Amazon, where an AI recruitment tool systematically disadvantaged female candidates. The system downgraded resumes containing indicators associated with women, reflecting biases present in the historical data on which it was trained.

Heterosexist bias is another type, and is defined as “conceptualizing human experience in strictly heterosexual terms and consequently ignoring, invalidating, or derogating homosexual behaviors and sexual orientation, and lesbian, gay male, and bisexual relationships and lifestyles” (Herek et al., 1991). Moreover, a study by Everly et al. (2016) shows that men are less likely to hire homosexual job candidates than heterosexual. On the other hand it was the opposite for women, which shows how recruiters bring their unconscious bias into the evaluation of candidates. It also highlights the need for a diverse work environment and recruiters with different genders, backgrounds, etc. (Everly et al., 2016).

Together, these studies show that different forms of algorithmic bias is a common occurrence, and it demonstrates that algorithmic systems are not inherently objective, but shaped by the data and assumptions underlying their design.

2.2.2 Algorithmic Fairness

Fairness is defined as “*the quality of treating people equally or in a way that is right or reasonable*” (Cambridge Dictionary, n.d.c).

From a broader perspective, fairness in AI systems is not solely a technical issue. Teo (2025) conceptualizes algorithmic fairness as a polycentric problem, where multiple objectives must be balanced simultaneously rather than optimized in isolation. These objectives include compliance with non-discrimination law, meaning equal treatment across individuals or groups, as well as system accuracy, referring to the ability of the AI system to produce correct

or reliable predictions. They also include the underlying purpose of the AI system, which may reflect organizational goals such as efficiency or performance.

According to Teo (2025), these elements are often in tension with one another, as efforts to improve one objective may negatively affect another. For instance, increasing system accuracy based on historical data may reinforce existing patterns that conflict with non-discrimination principles. At the same time, Wang et al. (2022) explains that adjustments made to reduce discriminatory outcomes may lower predictive performance. This illustrates that such trade-offs are inherent to the design and use of AI systems. This challenge is further reinforced by the existence of multiple definitions of algorithmic fairness, which are often incompatible and cannot be satisfied simultaneously.

These findings suggest that fairness cannot be achieved solely through technical adjustments to algorithms. Agbasiere and Nze-Igwe (2025) further demonstrate that ensuring fairness requires consideration across multiple stages of the AI system. Fairness is influenced not only by the training data, but also by how outcomes are defined and which variables are included in the model. In practice, decisions regarding what constitutes a “successful” candidate, as well as which attributes are considered relevant for evaluation, are often shaped by organizational priorities rather than fairness alone.

Research on fairness-aware machine learning highlights that different approaches to promoting fairness, such as reweighting data, applying fairness constraints, or adjusting model outputs, can produce varying outcomes depending on the context. This suggests that fairness is not a single measurable property, but rather depends on how it is operationalized within the system (Agbasiere & Nze-Igwe, 2025). In addition, the lack of comprehensive and universally applicable fairness metrics makes it difficult to consistently assess fairness across different contexts (Wang et al., 2022).

As a result, fairness in AI-driven recruitment should be understood as an ongoing and context-dependent process, rather than a fixed characteristic of an algorithm. Outcomes are shaped not only by technical design choices, but also by how recruiters interpret and act on algorithmic rankings, as well as what information is presented to them. Similarly, Möller and Petrovic (2025) find that candidates with relevant competencies are not always ranked at the top of lists. Respondents noted that factors such as resume layout and wording primarily influence outcomes, and the systems may also filter candidates based on their name, sex, education, or age. However, not all respondents experienced bias directly, but many acknowledged that it could occur.

This further illustrates that fairness in AI-driven recruitment is not only shaped by algorithmic design, but also by how systems are used in practice. Fabris et al. (2026) demonstrate that improving fairness within algorithms does not necessarily lead to fair outcomes in practice. Even when candidate rankings are adjusted to increase equal visibility, recruiters tend to focus on candidates at the top of the list (i.e., position bias). This means that candidates shown further down receive fewer opportunities, particularly when recruiters are under time constraints or when the interface does not encourage broader exploration. This highlights how

user behavior and system design can reinforce unequal outcomes, even when fairness has been addressed at the algorithmic level.

2.2.3 *Transparency*

A central issue in Human Resource Management is the lack of transparency in algorithmic decision-making, often described as the black box problem. This means the inability to understand or interpret how the algorithm's logic works when reaching specific decisions. AI-systems used in recruitment are often perceived as opaque, limiting insight into how the algorithms come to certain decisions (Möller & Petrovic, 2025). Möller and Petrovic (2025) found that employees usually cannot explain the reasoning behind the AI-generated decisions. They further explain that this is particularly problematic because it makes it difficult to explain the outcomes of the selection process, making it a major ethical challenge in AI-supported recruitment.

On one hand, limited transparency can reduce trust and lead to skepticism toward algorithmic decisions. On the other hand, research shows that individuals may still rely on algorithmic advice despite not fully understanding it, or reject it altogether depending on the context (Logg et al., 2019; Xiong & Kim, 2025).

One way to mitigate these issues is by increasing users' perceived controls. People are more likely to rely on algorithmic advice when they are given the ability to modify its outputs, even if it's in a limited way. This sense of control increases trust and the perceived effectiveness of the system (Dietvorst, 2016). It raises questions about whether perceived control reflects actual understanding or merely creates a sense of confidence without improving transparency.

The uncertainty is also reflected in recent research on algorithmic HRM, which highlights that professionals often experience a "fear of the unknown", due to a limited understanding of algorithmic processes. As a result, perceptions of algorithmic decision-making are often ambivalent, combining both skepticism and appreciation depending on the level of transparency and user understanding (Tandon et al., 2025).

These challenges also have important implications for accountability. Limited transparency makes it difficult to justify, interpret, and critically evaluate algorithmic decisions, which is an issue further discussed in the following section.

2.2.4 *Accountability*

Accountability addresses who is responsible for algorithmic decisions and their outcomes. The opacity of algorithmic systems raises concerns about accountability (Kumaresh et al., 2025; Möller & Petrovic, 2025). Ensuring accountability in recruitment contexts means that

there is a human decision-maker ultimately responsible for the outcomes, and it is essential for maintaining ethical standards (Yanamala, 2023).

However, algorithmic systems may create an illusion of control where recruiters overestimate the effectiveness and rely too heavily on it. This is closely connected to the concept of human sense-making, where decision-makers interpret and act upon algorithmic outputs. When AI systems shift from descriptive to prescriptive roles, there is a risk that recruiters rely less on their own reasoning or judgment. This can lead to what has been described as “nomolatriy”, where algorithmic outputs are treated as superior. The inability to question or challenge algorithmic reasoning can cause automation bias and diffused accountability, as responsibility shifts from humans to the systems (Karami & Ghasemi, 2026; Leicht-Deobald et al., 2019).

At the same time research shows that recruiters do not just follow algorithmic recommendations without question. Instead, they actively work with and sometimes challenge them. For example, recruiters may interpret recommendations differently or even ignore the system’s suggestions if they see limitations. This process, known as “AI crafting,” describes how users try to stay in control and make sure that algorithmic outputs match their own judgment and professional standards (Laukkarinen, 2025). Human oversight remains essential, since the responsibility ultimately lies with recruiters. These systems create the most value when used as decision support rather than decision replacement. Individuals who actively interpret and contextualize algorithmic recommendations are better able to integrate them into meaningful decision-making processes (Karami & Ghasemi, 2026).

2.3 Legal and regulatory considerations in recruitment

2.3.1 Transparency and regulatory challenges

As discussed in the previous sections, the use of AI in recruitment raises questions about transparency and accountability. Teo et al. (2025) show that AI systems are shaped by design choices and underlying assumptions, which means that their outputs are influenced by how they are built rather than being neutral. The study also highlights that these systems often rely on human-like forms of interaction. Because they imitate human communication, users usually see them as intuitive and trustworthy, and therefore they stop questioning the output.

At the same time, these systems are complex and lack transparency, making it difficult to understand how specific outcomes are produced. Burrell (2016) argues that machine learning systems create opacity because their decision-making processes are based on learned representations derived from large amounts of training data, rather than clearly interpretable rules. As a result, the underlying reasoning behind algorithmic outputs may become difficult for humans to critically evaluate. This lack of transparency creates challenges in assessing

how algorithmic recommendations are produced and whether they can be considered reliable or fair.

Research on AI also emphasizes that these systems should be understood in relation to the social and organizational contexts in which they are used. Their integration into society raises broader ethical concerns regarding the consequences they may have for individuals, organizations, and existing social structures. A socio-technical perspective highlights how AI systems interact with human decision-making and institutional structures, which may influence how outcomes are interpreted and acted upon (Mbiazi et al., 2025)

Goodman and Flaxman (2017) highlight the “right to explanation” as an important issue in AI-driven decision-making, particularly in relation to the GDPR. The principle implies that individuals should be able to understand the reasoning behind decisions that affect them. However, Wachter et al. (2016), argue that the GDPR does not establish a clear right to a full explanation, leaving uncertainty around what level of transparency is required. The complexity of machine learning systems adds to this, as explanations are often difficult to provide due to both technical limitations and system design (Burrell, 2016). Organizations are therefore left to interpret what counts as a sufficient explanation.

From a legal perspective, the use of AI in recruitment must comply with data protection and non-discrimination requirements. Article 22 of the General Data Protection Regulation (GDPR) states that individuals have the right to not be subject to decisions based solely on automated processing, including profiling, when such decisions have legal or similarly significant effects (European Union, 2016). This may include automated processes such as filtering or ranking of candidates.

Regulatory initiatives such as the EU AI Act aim to mitigate these risks, by introducing requirements for how AI systems should be developed and used (European Parliamentary Research Service, 2024). In Europe, the GDPR also regulates how personal data is processed, including restrictions related to use of sensitive data in automated decision-making (Veale & Zuiderveen Borgesius, 2021).

2.3.2 Accountability and human oversight

Leicht-Deobald et al. (2019) mean that the use of algorithmic decision-making in HR may reduce human sense-making, and encourage increased reliance on automated processes. This raises concern regarding responsibility and ethical oversight. Similarly, Karami and Ghasemi (2026) argue that accountability may become more difficult to maintain when HR-leaders adopt algorithmic recommendations, without critically questioning them. This diffusion of responsibility results in challenges both in terms of legal accountability and in ensuring that decisions can be justified.

The use of AI in decision-making raises questions about responsibility, particularly as outcomes are shaped through interactions between human judgment and AI systems (Floridi et al., 2018). It becomes problematic in recruitment as decisions must often be justified, and

can have direct consequences for individuals. At the same time, responsibility is not always clearly defined (Raji et al., 2020).

However, regulating AI remains challenging. Existing legal frameworks do not always reflect how AI systems function in practice (Teo et al., 2025), particularly as these systems are context-dependent (Veale & Zuiderveen Borgesius, 2021). Therefore, a gap emerges between formal regulation and practical implementation, requiring organizations to interpret and apply legal requirements in complex real-world settings (Mbiazi et al., 2025).

2.4 Mitigating Ethical Risks in AI-Driven Recruitment

2.4.1 Transparency and Explainability Measures

According to Yanamala (2023), transparency in algorithmic decision-making is considered a key factor in mitigating ethical risks in AI-driven recruitment. This involves improving understanding of how AI systems generate outcomes, and ensuring that relevant stakeholders are informed about the criteria underlying these decisions. Greater transparency enhances accountability, facilitates the identification of potential biases, and contributes to building trust in AI-supported decision-making processes. Hunkenschroer and Kriebitz (2022) suggests that organizations can reduce the complexity of their AI systems and its algorithms, so managers understand and can explain it.

2.4.2 Bias Detection and Fairness-aware Models

According to Yanamala (2023), HR managers should make sure that the datasets used to train the AI models are representative of diverse populations, to mitigate the risk of biased outcomes related to gender, race, or ethnicity. However, existing literature emphasizes that data diversity alone is insufficient to fully address bias. Bias detection mechanisms, such as fairness metrics and regular audits, are necessary to identify discriminatory patterns in model outputs. In addition, fairness-aware models can be applied to actively mitigate bias during the decision-making process. This can be done by for example adjusting model parameters, or introducing constraints aimed at promoting equitable outcomes across demographic groups. These approaches illustrate that mitigating bias in AI recruitment requires both improved data practices and technical interventions within the modeling process.

Hunkenschroer and Kriebitz (2022) explains a bias-mitigation technique where HR managers can remove sensitive or potentially biased-inducing information about a candidate. This includes excluding personal or social information, making sure it is not used during

evaluation. Yanamala (2023) also explains that regular auditing of AI systems plays a key role in uncovering and mitigating bias. This process involves analyzing both the data used to train models and their resulting outputs to identify discriminatory patterns and promote equitable decision-making.

2.4.3 Human Oversight in Decision-making

Incorporating human oversight at critical stages of the decision-making process allows automated outputs to be evaluated for fairness. In recruitment contexts, this means that AI-generated recommendations are reviewed by human decision-makers before final hiring decisions are made. In turn it reduces the risk of biased or unjust outcomes (Yanamala, 2023).

Hunkenschroer and Kriebitz (2022) state that *“hiring decisions should always be made by an AI-informed human rather than by AI alone”*. While AI can assist in evaluating candidates, final hiring decisions should remain under human control, ensuring that algorithmic outputs are interpreted and validated by humans.

2.4.4 Governance, Policies, and Regulatory Compliance

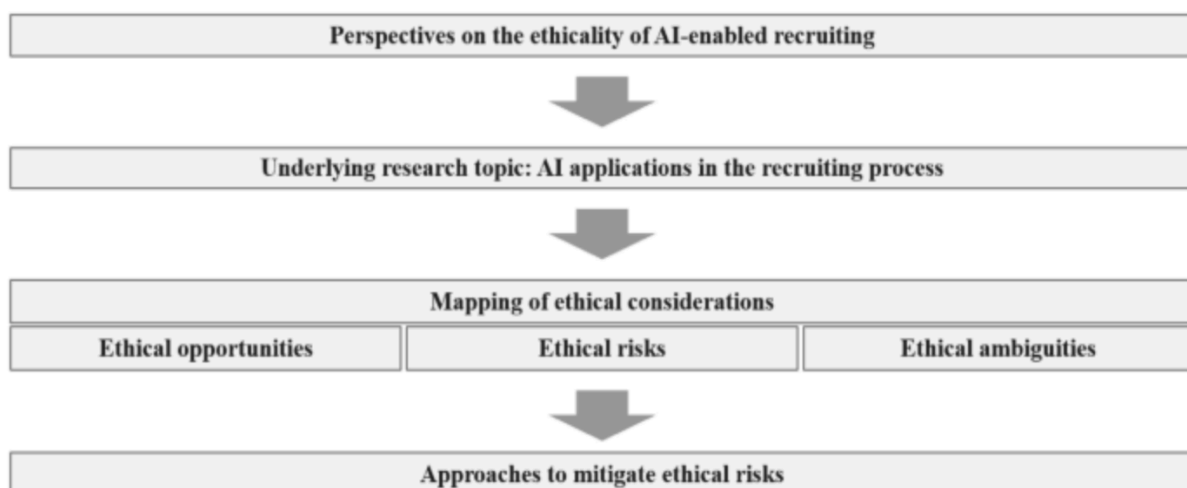
In order to mitigate the ethical risks in governance, policies, and regulatory compliance, Yanamala (2023) suggests limiting the amount of personal data exposed to the AI systems. Another important aspect is clearly informing the candidate's consent. Companies should also have strict internal policies, and adhere to regulatory compliance to mitigate the risks. Hunkenschroer and Kriebitz (2022) talk about the practice of data minimization, where organizations should only collect data classified as necessary.

3 Theory

3.1 Ethical Framework for AI-enabled recruiting

The Ethics of AI-Enabled Recruiting framework, developed by Hunkenschroer and Luetge (2022), is a systematic review-based framework that maps the ethical dimensions of AI use across the whole recruitment process. As aforementioned, these different stages include outreach, screening, assessment, and facilitation. The framework was developed through a structured literature review of 51 articles and synthesizes ethical considerations into three categories: opportunities, risks, and ambiguities, as illustrated in Figure 1. The framework is not limited to a single perspective, but combines contributions from theoretical, practitioner, legal, technical, and descriptive research. Based on these perspectives, Hunkenschroer and Luetge (2022) identified several ethical themes related to AI-enabled recruitment.

Figure 1. Ethical Framework of AI-Enabled Recruiting (Hunkenschroer & Luetge, 2022)



3.2 Application of the Ethical Framework

The study applies the ethical framework developed by Hunkenschroer and Luetge (2022), using the same categories of ethical opportunities, risks, and ambiguities, as the analytical structure. However, the specific themes presented in the original framework were refined and reorganized to better reflect the empirical findings of this study, which is further discussed in section 3.2.1. Table 1 (adapted from Hunkenschroer & Luetge, 2025) shows the updated themes.

The study aims to explore and develop an understanding of the ethical implications of using AI in recruitment, making the framework by Hunkenschroer and Luetge particularly suitable. The framework addresses ethical dimensions of AI-enabled recruitment across multiple stages of the recruitment process, which makes it broadly applicable across different empirical contexts. This research is limited to the stages of the recruitment process where the respondents use AI. In addition, the categorization of ethical opportunities, risks, and ambiguities, supports a nuanced analysis of AI-supported recruitment. It allows the study to examine both beneficial and problematic aspects of AI use rather than assuming a solely positive or negative perspective. The framework also aligns well with the qualitative approach to the interviews, as its predefined categories support a systematic, yet flexible, interpretation of HR professionals' experiences and perspectives on AI-supported recruitment.

Table 1. (adapted from Hunkenschroer & Luetge, 2025) Ethical considerations

Ethical opportunities	Ethical risks	Ethical ambiguities
Efficiency and process support	Introduction to algorithmic bias	Effect on workforce diversity
Reduction of human bias	Lack of transparency	Impact on assessment validity and accuracy
	Potential loss of human oversight and accountability	Perceived fairness
	Legal and ethical risk to data usage	The interpretation and application of AI regulation
	Assessment limitations and misinterpretation	

3.2.1 Refinement of Framework Themes

The empirical material was initially structured according to the categories of ethical opportunities, risks, and ambiguities. However, certain derived themes were adjusted and combined in order to better capture patterns. For example, the themes related to potential loss of human oversight, and diffused accountability, were grouped under ethical risks.

During the analysis, some themes were also found to overlap conceptually. In the discussion section these overlaps enabled broader analytical interpretations across categories, which resulted in the removal of the subsection “Impact on assessment validity and accuracy”.

Themes from the framework that were not identified in the empirical material, were not included in the empirical chapter or further discussed in the analysis. At the same time, the

empirical findings revealed additional themes that were not explicitly covered within the original framework. Consequently, three additional subsections were introduced; “Legal and ethical risks related to data usage” and “Assessment limitations and misinterpretation” under ethical risks, and “The interpretation and application of AI regulation” under ethical ambiguities. These themes were considered relevant to the study of ethical considerations because they involve the handling of sensitive candidate data, and challenges related to privacy and transparency. “Assessment limitations and misinterpretation” includes experiences and findings related to how AI-supported evaluations may lead to unfair or inconsistent outcomes and are therefore associated with ethical risks.

4 Methodology

4.1 Qualitative research

Qualitative research is a methodology that is used to achieve a more profound and detailed understanding of specific problems, events, or situations (Björklund & Paulsson, 2012). This research requires a qualitative approach since the purpose was to explore and develop an understanding of what ethical challenges emerge in AI-supported recruitment, and how these are experienced and managed by HR professionals, in practice. Qualitative research is suitable for investigating complex issues, as it focuses on individuals' perceptions and experiences (Kalu & Bwalya, 2017).

The research question requires insights into how HR professionals think and act when working with AI-supported tools in the recruitment process. Ethical issues such as fairness, bias, transparency and accountability are not only of a technical nature, but they are also shaped by human judgment and the organizational context. These aspects are difficult to measure with numbers, and cannot be fully captured with quantitative methods as they require human experiences, to gain a deeper understanding. A qualitative approach enables a more nuanced understanding of how HR professionals reason about, and manage, ethical considerations in AI-supported recruitment. This approach also allows respondents to elaborate on their experiences and provide detailed explanations. According to Oates et al. (2022), a qualitative approach allows for rich and detailed information that a quantitative approach cannot provide. With a qualitative approach, this research could better explore the collaboration between humans and AI, and how ethical issues are addressed in practice.

4.2 Semi-structured interviews

For this research semi-structured interviews were conducted, meaning that the themes for the interviews were predetermined. However, the wording of the questions were adjusted during the interviews. This is done in accordance with the respondents' answers and reactions to previously asked questions (Björklund & Paulsson, 2012).

Semi-structured interviews were the most suitable approach because the participants used AI systems in different ways, and had different levels of experience. Therefore, it was not possible to follow a standardized set of questions throughout all interviews. Instead, the interviews were adapted to each participant by adjusting the wording, order, and focus of the questions. According to Yin (2016), such adaptability is a strength in qualitative interviews, since researchers can explore relevant issues in greater depth and capture differences in experiences across participants.

4.2.1 Sampling approach

For this study, non-probability sampling techniques were employed. This means that participants are selected based on non-random criteria. Specifically, the study used an instant purposive sampling, used when research requires specific information and selects participants purposely based on characteristics (Bryman, 2018). In this case it concerned relevant knowledge or experience related to the research topic. Participants for this study had either worked with AI-based systems, had relevant experience and understanding of such technologies, or could explain why they did not decide to use such technologies. All participants were screened to confirm that they met the inclusion criteria, and only participants who showed familiarity with the topic were included in this study.

Participants were recruited through extended personal networks. In practice, this means that initial contact was facilitated through people with broader professional networks, such as family members, who helped connect the researchers with relevant HR professionals. However, the researchers did not have prior personal relationships with any of the respondents themselves. The approach was used as a practical and time-efficient means of gaining access to suitable participants. While this may introduce a potential risk of bias, there is no indication that it influenced the respondents' answers. The participants were selected based on their professional expertise, and were encouraged to respond freely and independently.

All respondents are anonymous. Additional identifying information is only included when participants explicitly stated that they did not mind being identified. For those who wished to remain fully anonymous, only their professional role is presented (e.g., Talent Acquisition Partner) to demonstrate their relevance to the study. It ensures that they cannot be traced through additional details (e.g., name, company, or location). A summary of the respondents is provided in Table 2.

Table 2. Summary of Respondents

Respondent	Role / Experience	Language	Interview time (min : sec)
Respondent 1 (R1)	CEO at a talent acquisition consulting and recruitment company called “Zoey”	Swedish	25:14
Respondent 2 (R2)	Talent Acquisition Partner at “Sambla Group”	Swedish	21:12
Respondent 3 (R3)	Talent Acquisition Partner	English	29:40
Respondent 4 (R4)	Talent Advisory Group	English	45:46
Respondent 5 (R5)	Recruiter at “Stegra”, responsible for the whole	Swedish	-

	recruitment process		
Respondent 6 (R6)	COO and People Advisor at a consultant company called “Recommended By”	Swedish	26:40

4.2.2 Design of interview guide

The interview guide was developed as a semi-structured tool, meaning that it was designed to guide the conversation rather than function as a fixed script. This approach allows the researcher to ensure that key topics are covered while still giving participants the opportunity to elaborate on their experiences and introduce new perspectives during the interview (Oates et al., 2022). As aforementioned, the development of the interview guide was grounded in the framework *Ethics of AI-Enabled Recruiting* (Hunkenschroer & Luetge, 2022). This means that the categories of ethical opportunities, risks, and ambiguities guided the development of the interview questions.

Each theme was translated into open-ended questions. At the same time, several themes could be addressed within a single question, as semi-structured interviews allow for flexibility and enable multiple aspects to be explored within the same discussion (Bryman, 2018). This allowed participants to describe their experiences in their own words. Developing the interview questions based on a theoretical framework helped ensure that the data collection was conceptually grounded and aligned with the research focus.

According to Bryman (2018) the interview questions should follow a predetermined order, but only to a certain extent, since the follow-up questions can change the direction of the interview. In line with this, the interview guide was organized around a logical progression of topics. The interviews began with introductory questions about the participants’ roles and experiences to establish context and create a comfortable setting. The discussion then moved to ethical opportunities associated with AI in recruitment, and then addressing ethical risks and how these are managed by HR professionals. This structure facilitated a natural flow in the conversation and allowed more sensitive or complex topics to be introduced progressively. See Appendix B for a general interview guide.

4.2.3 Interview preparation

It is important to be well-prepared about a topic when conducting an interview (Bryman, 2018). Therefore, all interviewers participated in the different parts of the literature review, and ensured that they were familiar with everything. In addition, all interviewers were well informed about the theoretical framework used in the study, to ensure consistency in how questions were asked, and how the topic was approached. Bryman (2018) also emphasizes

that interviewers need to be structured, while also encouraging new perspectives and ideas. It is important to ask relevant follow-up questions and steer the interview so it will provide valuable information for the analysis, without asking leading questions. This was kept in mind during the whole data collection procedure.

4.2.4 Data collection procedure

The interviews were conducted digitally via either Google Meet or Teams, which allowed flexibility in scheduling. Each interview lasted approximately 20 to 45 minutes (see Table 2 above). However, due to unforeseen circumstances, Respondent 5 was unable to participate in a live interview. Instead, the interview questions were sent to the respondent, who provided written responses. The responses were analyzed using the same coding procedure as the transcribed interviews, which is explained in detail in section “4.4.2 Thematic analysis”.

During the interviews, there were always two members of the research team present. Due to the nature of semi-structured interviews (i.e., based on themes rather than a fixed set of questions), interviewers must continuously adapt, probe, and interpret responses in real time (Bryman, 2018). Having two researchers present was an active decision that contributed to a more valuable data collection process. It allowed for a division of attention, where one researcher led the conversation following the interview guide, while the other observed and identified additional points for probing. This increased the likelihood of capturing richer and more nuanced data, as different researchers may notice different aspects of participants’ responses. It also reduced the risk of missing relevant information and supported more immediate reflection during the interview process.

Before the interviews, participants were informed about the purpose of the study, and their consent to participate and be recorded was obtained. Participants were also given the option to remain anonymous. Ethical considerations taken into account during the data collection process is further explained in section “4.4.1 Ethical considerations”. All interviews were later transcribed, which is explained in detail in section “4.3.1 Transcription”.

4.3 Analysis

4.3.1 Transcription

For the transcription of the interviews, the AI tool Whisper was used. This choice was based on recommendations from the Informatics Department at Lund University, primarily due to its data security advantages. Whisper can be run locally on a computer, meaning that neither audio nor video data needs to be uploaded to external servers. As a result, no sensitive

information is stored in the cloud, which reduces the risk of unauthorized access and strengthens the protection of respondents' privacy and integrity.

The interviews were conducted in both Swedish and English. The interviews were first transcribed in their original language using Whisper. Since the thesis is written in English, all quotes or excerpts included in the empirical findings (i.e., 5. *Empirical data*) are in English, meaning that Swedish quotes or experts in that section have been translated manually by the researchers. This approach was chosen to preserve the original meaning as closely as possible. However, despite these efforts, it is important to note that minor nuances or contextual meanings may have been lost in the translation process.

Following the initial transcription, all recordings were manually reviewed, and the transcripts were checked and corrected where necessary before analysis to ensure a high level of accuracy. While the majority of the transcriptions generated by Whisper were accurate, some adjustments were required. The following example is an output generated by Whisper. As the original interview was conducted in Swedish, the excerpt has been translated into English as the following:

“Exactly. You have an ATS-system, African Tracking System. When you apply for a job, you may for example use Team Taylor”. This is what the respondent actually said: *“Exactly. You have an ATS-system, Applicant Tracking System. When you apply for a job, you may for example use TeamTailor”*

This illustrates that minor transcription errors, such as incorrect word recognition, needed to be corrected before analysis. Moreover, some parts had to be completely omitted. For example, this was output generated by Whisper from another respondent:

“I'm one of the most innumerable divinhas in that uh what the business becomes in terms of online coaching and becoming introduced whenever i get forced and i'm having a hard time there's all sorts of things i need to accomplish i need to be interactively better in terms of how i canomic im очень comfortable getting to startelen to be actually if i can back up workflow, and AI can be used to help them on every step.”

This was not said by the respondent at all, which illustrates that Whisper occasionally produced incoherent or incorrect text, particularly in informal or unstructured parts of the recording.

As aforementioned, only two researchers were present during each interview. Therefore, the researcher who did not participate in a given interview was responsible for transcribing it, in order to ensure familiarity with the data which was ultimately important for the analysis. Transcribed interviews can be found in Appendix C.

4.3.2 Thematic analysis

Thematic analysis is a well-known method to identify and analyze patterns or themes using categories when analyzing qualitative data. The method involved creating categories that constitute codes that can be read out or discovered in the material. Thematic analysis enables a detailed and in-depth analysis of specific aspects of the collected data (Braun & Clarke, 2006). This method is considered suitable for the study, as it allows the empirical material to be structured around recurring themes related to ethical considerations. Given the exploratory nature of the research question, a flexible analytical approach was needed to capture how such considerations emerge in practice rather than testing predefined relationships. Thematic analysis allows for this flexibility, as it makes it possible to identify patterns across participants' experiences, while still remaining open to unexpected findings (Braun & Clarke, 2006).

The analysis was guided by the theoretical framework presented in Chapter 3, where ethical opportunities, risk, and ambiguities constituted the main analytical categories. Within each category, specific themes derived from the framework were used as codes in the analysis. In this sense, the study followed a primarily deductive approach. At the same time, the researchers remained open to additional patterns and nuances that emerged from the interview material.

This approach aligns with Bryman's (2018) description of the so-called "*framework method*", which is a matrix-based strategy for organizing and analyzing qualitative data. In this study, the empirical material was structured in a spreadsheet in which interview excerpts were categorized according to themes and ethical considerations. This enabled systematic comparisons across respondents and analytical categories.

The coding process began with all researchers reading through the transcripts to become familiar with the data. Relevant statements were then coded according to the predefined themes. These codes were used to identify similarities and differences across respondents and to structure the empirical findings. To support consistency in the analysis, the coding was documented in an Google Sheet coding scheme. Each respondent was assigned a number from R1 to R6, and the themes were color-coded to make the analysis more systematic and transparent. The coding scheme is presented in Table 3.

Table 3. Coding scheme for data analysis

Main theme	Colour	Code	Description
Bias	RED	BS	Statements about different types of bias that may occur in AI-supported recruitment.
Fairness	PURPLE	FS	Statements about how fairness is perceived, evaluated or addressed.
Transparency	BLUE	TY	Statements about how understandable,

			explainable, or visible the AI systems are.
Accountability	GREEN	AY	Statements about who is responsible for AI-supported recruitment decisions.
Human–AI interaction	YELLOW	HA	Statements about how recruiters use, question, override or rely on AI outputs.
Efficiency	PINK	EY	Statements about time-saving, workload reduction, or handling many applications.
Consistency	CYAN	CY	Statements about standardizing processes or consistent output.
Legality	ORANGE	LY	Statements about GDPR, AI-act, and other compliance and regulations.

Table 4 presents an example of how the empirical material was coded and categorized within the analytical structure.

Table 4. Extract from Coding Scheme

Row	Respondent	Quote	Code	Ethical consideration
30	R2	“there are also advantages from a discrimination perspective; the process can become more objective and reduce the risk of me unconsciously selecting someone based on similarities [...] All candidates are then given the same conditions”	BS, FS	Opportunity

Braune and Clarke (2006) emphasize, the method also means that the researchers’ assumptions and interpretations in the analysis have a significant impact, which is therefore evident in the analysis and discussion material. Therefore, the coding and interpretation were discussed among the researchers to ensure consistency, as further elaborated in the next section.

4.4 Reliability & validity

In qualitative research, reliability and validity are used to assess the trustworthiness of a study, although they are understood differently compared to quantitative research. Reliability refers to the consistency of the research process, while validity concerns the credibility and accuracy of the findings (Creswell, 2009)

To ensure reliability, the research process was conducted in a systematic and transparent manner. All interviews were recorded and transcribed, allowing the analysis to be based on consistent data. The material was reviewed multiple times during the coding and interpretation process, which contributed to maintaining consistency in how the data was understood. Since all researchers were involved in the analysis, interpretations could be discussed and compared, reducing the risk of individual bias and strengthening the overall consistency of the findings.

Validity was addressed through the design of the study and the data collection process. The interview guide was developed based on relevant literature and aligned with the research question, ensuring that the data collected was relevant to the purpose of the study. During the interviews, follow-up questions were used to clarify responses and gain deeper insights into the participants' perspectives. Providing detailed descriptions of the findings also contributes to strengthening validity, as it allows for a more accurate and nuanced representation of the empirical material (Creswell, 2009).

4.4.1 Ethical considerations

Ethical considerations were taken into account throughout the research process in order to create a safe and open setting for the participants. Before the interviews, participants were informed about the purpose of the study and how the collected data would be used. They were also made aware that participation was voluntary, and that they could withdraw at any time (Oates et al., 2022).

Another aspect to this was to carefully maintain confidentiality. Participants were given the option to remain anonymous, and no identifiable information has been included in the thesis for respondents who chose anonymity. All data, including interview recordings and transcripts, was stored securely to prevent unauthorized access. This is especially important in qualitative research, where participants may share personal experiences or opinions that they expect to remain private (Oates et al., 2022).

The responsibility of the researchers was also considered. The data has been reported accurately and without manipulation, and care has been taken to present the findings in a transparent and honest manner. This includes ensuring that the analysis reflects the participants' perspectives rather than being influenced by the researchers' own assumptions (Oates et al., 2022).

5 Empirical data

5.1 Ethical opportunities

5.1.1 Efficiency and process support

Respondents describe how AI contributes to more structured and standardized workflows, reducing variation in how recruitment tasks are performed (R1; R3; R4; R5). This is particularly visible in communication and coordination processes, where AI supports more consistent interactions with candidates. For example, one respondent notes that *“we become faster and more precise in our communication”* (R1, 10), while another highlights improvements in *“candidate communication (emails, feedback)”* (R5, 186).

Improved communication is also linked to a more reliable candidate experience. Automation enables timely and standardized interactions, reducing the risk of human error and ensuring that candidates are not overlooked. As one respondent explains, *“All candidates who reach a certain stage automatically receive an email, which reduces the risk of human error [...] the process becomes more secure and consistent”* (R2, 222). In addition, automated responses allow candidates to receive feedback regardless of when they apply (R4).

Efficiency gains are closely connected to these improvements. AI is described as enabling faster processes, reducing administrative workload, and supporting the handling of large volumes of candidates (R1; R2; R3; R4). It is primarily used for administrative and repetitive tasks, such as handling candidate data, communication, and coordination of recruitment processes. Several respondents describe AI as a support function that reduces manual effort and streamlines workflows, for example through *“automated steps for administrative tasks that do not require any expertise”* (R2, 24) and *“an administrative support tool and a project management support tool”* (R1, 9). This is also reflected in reduced repetition, as one respondent notes that *“it still reduces the going back and forth on different profiles”* (R3, 50).

The importance of efficiency becomes particularly evident in contexts with large applicant volumes, where manual screening is no longer feasible. As one respondent states, *“there's no company that has enough recruiters to look at 18,000 applicants for one job”* (R4, 114). Similarly, another respondent describes how time constraints shape selection practices, as recruiters often prioritize recently submitted applications due to limited capacity (R3). However they argue that AI in this case would make it more fair to the applicants by at least looking at all of them, while also offering cost-related advantages, as it is described as *“cheaper than even the cheapest human resource”* (R4, 131).

5.1.2 Reduction of human bias

The reduction of human bias is described as a potential benefit of using AI in recruitment. Several respondents highlight that human decision-making is subjective, which may influence candidate evaluations. In this context, AI is perceived as a tool that can contribute to more objective and standardized assessments (R2; R4). Respondents describe how standardized processes may reduce the risk of unconscious bias, particularly in situations where recruiters might otherwise rely on personal preferences or perceived similarities. As R2 states *“there are also advantages from a discrimination perspective; the process can become more objective and reduce the risk of me unconsciously selecting someone based on similarities [...] All candidates are then given the same conditions”* (R2, 30).

AI is also described as a complementary tool that can support recruiters in identifying and reflecting on their own biases. Respondent 4 highlights this by explaining that AI can function as an additional layer of evaluation, stating that *“AI can help you be more inclusive and fair because of human unconscious bias, and using AI specifically as a second opinion to actually be like a watchdog, to look out for those types of things”* (R4, 128). The respondent further sees future opportunities, by describing *“I do think we’re going to reach a point where AI can be consistently less biased than people”* (R4, 112).

5.2 Ethical risks

5.2.1 Introduction to algorithmic bias

Concerns about algorithmic bias are recurring across respondents, particularly in relation to how AI systems rely on historical data and predefined patterns. Several respondents highlight the risk that such systems may reproduce past recruitment outcomes, rather than evaluate candidates based on current requirements (R1; R3; R4). This is especially evident in discussions of diversity, where R1 notes a concern that AI may favor profiles that have previously been considered successful, leading to homogeneous outcomes. The respondent highlights this by explaining that *“we want to stand for diversity, and we are afraid that AI goes off of what has been most successful in the past. So that we end up only getting, like, 50-year-old white men for our roles, and that’s not what we want”* (R1, 4).

Similar concerns are raised regarding embedded bias in training data, where AI systems may inherit existing patterns. This means that AI systems may repeat earlier hiring patterns, instead of evaluating candidates independently. As one respondent explains, there is a risk that bias becomes *“encoded or embedded”* in the system (R4, 107). In addition, respondents highlight how indirect or proxy variables, such as names or addresses, may influence outcomes. For example, one respondent notes that *“certain addresses can be proxies for socioeconomic status, which can then be a proxy for race and ethnicity”* (R4, 109), while another points to concerns about how names may affect evaluation (R3, 73). This suggests that bias may arise even when sensitive attributes are not explicitly included.

5.2.2 Lack of transparency

Respondents describe limited insight into how algorithmic outputs are generated, making it difficult to understand or interpret recommendations (R1; R2; R3; R6). Understanding is often based on general assumptions, rather than concrete knowledge. For example, R1 states *“it is perhaps more of a basic understanding of how the algorithms work, or a general idea that they rely on what has been successful in the past”* (R1, 13). This is reflected in respondents’ descriptions of technical understanding. As an example, one respondent states *“I personally do not have very deep technical expertise”* (R6, 210), while another notes that *“the buyers don’t really fully understand how they work”* (R4, 142).

Uncertainty about how decisions are produced is also evident in how respondents question AI-generated outcomes. For instance, one respondent asks *“AI selecting the top 10 candidates. Based on what?”* (R3, 47). Such statements indicate limited explainability, where the reasoning behind rankings or recommendations remains unclear. The same respondent states *“I don’t know why they selected this person. I never would have.”* (R3, 68). The quote reflects uncertainty regarding the AI system’s ability to accurately interpret and assess candidate suitability.

5.2.3 Potential loss of human oversight and accountability

All respondents highlight the risk of potential loss of human oversight, and choose to have a human making important decisions (R1; R2; R3; R4; R5; R6). For instance, R1 states *“we believe that the selection process needs to involve a human”* (R1, 14).

However, R4 states that human oversight may become weakened when AI is introduced into the recruitment process. The respondent describes situations where time pressure or high workloads increase the risk of relying too heavily on system outputs, rather than critically evaluating them. The respondent expresses this as *“you might have some users that are under pressure”* (R4, 100).

R4 means that there is a tendency to accept AI-generated results without thorough review. For example, *“they can’t fall asleep at the wheel, so to speak, and just kind of trust it”* (R4, 102). The respondent further highlights how involvement can become superficial, by explaining *“if I’m in the loop, but I’m just saying ‘it looks good’ because I’m quickly scanning the results, I’m not actually looking at each profile”* (R4, 137).

Another issue emphasized by R4 is the risk of *“rubber-stamping,”* where human oversight becomes superficial if users simply accept AI-generated recommendations, without critical evaluation. Moreover, even when systems are designed to support decision-making rather than replace it, and when clear guidelines are in place, human involvement may still be limited in practice. As R4 explains *“even when you tell people to do things, they may not necessarily actually do that”* (R4, 140).

5.2.4 Legal and ethical risks related to data usage

Respondents point to limited understanding of how data is handled across different actors, within organizations. For example, one respondent states *“when it comes to the AI Act, we more or less have a general understanding of it”* (R1, 220), and another respondent expresses that *“it made me realize that i do not have full insight into how my colleagues handle and store such information”* (R2, 38).

It is also mentioned that candidates may not be fully aware of how their data is used in AI-supported processes. As one respondent notes, *“but from the candidate’s perspective, I do not think it is very noticeable that we use AI in the background”* (R2, 37).

Respondents also describe practical challenges related to the handling of personal data in AI tools. For example, one respondent raises concern regarding how sensitive information is managed when using external systems, *“for example, looking at how we ensure that we do not input personal data into ChatGPT. in that regard, I have been more concerned with making sure that we use the paid version, which they say is supposed to be more secure. But I do not have any kind of follow-up control regarding that”* (R1, 18).

5.2.5 Assessment limitations and misinterpretation

During interviews, a new theme emerged as several respondents brought up that AI systems may misinterpret or overlook relevant information, particularly when candidates present similar experiences in different ways. Respondents describe how variations in job titles or CV formats can lead to qualified candidates being excluded (R2; R3; R4; R6). For example, R4 says *“unless the AI [...] understands that a business analyst might have these 17 other titles [...] you will be excluding and missing people that are qualified”* (R4, 104). Another respondent says *“A person can describe their experience in very different ways while still having similar experience. So it is precisely those subtle differences that I think AI may not be able to capture properly yet”* (R6, 194).

Concerns regarding accuracy are further emphasized through examples of incorrect outputs by the system. R2 explained how an AI-generated summary contained incorrect information, explaining that *“it created a summary and said that the person had not worked with python, even though they had done so for five months [...] there were serious errors that made me hesitant to use it because of the risk of losing candidates and making incorrect assessments”* (R2, 22).

Respondents highlight that AI systems may fail to capture the full scope of candidate information, when processing CVs. They state that relevant details may be overlooked rather than explicitly misinterpreted. As one respondent explains *“the ones that scan an entire CV, I*

think they miss quite a lot of things” (R6, 199). Similarly, another respondent points to how formatting can influence outcomes, stating *“we could be missing candidates that maybe don't have the right formatting”* (R3, 61).

5.3 Ethical ambiguities

5.3.1 Effect on workforce diversity

There are differing perspectives on how AI affects workforce diversity. On one hand, AI is described as having the potential to support more objective and inclusive recruitment by applying consistent criteria across candidates (R2; R4). On the other hand, respondents express concern that AI may reinforce existing patterns and reduce diversity by relying on historical data or predefined structures (R1; R3; R4)

This creates tension between standardization and bias. While consistency in evaluation may reduce individual subjectivity, it may also reproduce structural biases embedded in past recruitment decisions. As one respondent explains, *“humans have unconscious bias”* (R4, 111), suggesting that bias is not limited to AI systems, but is also present in human decision-making.

This tension becomes more complex when considering the interaction between AI and human judgment. Respondents emphasize that AI does not operate independently but is embedded within a broader recruitment process where human involvement remains central. For example, one respondent notes that *“even if you had a very low-bias AI process, as soon as you get to the human interview, you reinsert the human bias”* (R4, 120). However, all respondents highlight the importance of maintaining human interaction as a deliberate choice in the process (R1; R2; R3; R4; R5; R6).

These perspectives illustrate that the challenge is not simply whether bias exists, but how it shifts between AI and human actors. As one respondent reflects, *“zero bias is probably impossible [...] but if you can create a system where bias is equalized, then it becomes fair”* (R4, 133). This suggests that fairness is not understood as the elimination of bias, but as a balance between different sources of bias.

5.3.2 Impact on assessment validity and accuracy

AI-generated outputs are often used to support the evaluation of candidates, but there is ambiguity regarding what these outputs actually represent and how they influence decision-making in practice.

Although a produced list is often perceived as a neutral way of organizing candidates, respondents highlight that it may implicitly function as a form of recommendation. As one respondent explains, even without explicit scoring, *“You could argue, however, that any matching solution is already making recommendations because it's an ordered list. So even if there are no scores involved, it's almost like using Google, right? Most people think that the top results beyond the sponsored are the best”* (R4, 93).

This creates uncertainty regarding what ranking actually represents. While AI may not explicitly score candidates, the presentation of results as an ordered list can influence how they are interpreted and acted upon. From a human perspective, higher-ranked candidates are more likely to receive attention, suggesting that ranking itself plays an active role in the assessment process. R4 states: *“I mean, back when it's just keyword frequency, that's an algorithm. That's ranking people. You could argue that that's essentially making some recommendation to people, that people typically process top down [...]”* (R4, 194)

5.3.3 Perceived fairness

Ambiguity also emerges in how candidate qualifications are represented and interpreted in AI-supported screening. Respondents highlight that AI systems rely heavily on the specific words used in CVs and applications, even though candidates may describe similar experiences and qualifications in different ways (R2; R3; R4; R6). This makes it difficult for AI systems to consistently identify and evaluate competence. As one respondent states, *“people are more than just the words that they decide to use”* (R4, 103). This creates uncertainty in how candidates are identified and evaluated, as relevant profiles may not be recognized if they do not match expected terminology. The same respondent explains that unless the system accounts for variations in job titles, *“you will be excluding and missing people that are qualified”* (R4, 104). The ambiguity therefore lies in the gap between how competence is expressed by candidates, how it is interpreted by the AI system, and how it is understood by recruiters.

Another respondent describes AI-supported recruitment as *“there's a little bit of give and take”* (R3, 62). This suggests that fairness in recruitment is not perceived as absolute, but rather as something that may involve compromises, and trade-offs, in practice. The findings therefore indicate an ambiguity in how fairness is understood, where organizations may accept certain limitations in fairness while still perceiving the overall recruitment process as sufficiently fair.

5.3.4 The interpretation and application of AI regulation

While regulations such as the GDPR establish boundaries for automated decision-making, respondents highlight that these rules are often formulated in general terms, leaving room for interpretation. *“there's a lot of grey area”* (R4, 123) which requires organizations to determine

for themselves what constitutes compliant use. R4 also explains that *“And then sometimes legal ends up driving everything. And so the business actually doesn't get the benefit of AI. Because they are not even using a fraction of what they're allowed to use”* (R4, 148).

This creates a tension between compliance and utilization. In navigating this uncertainty, some organizations adopt a cautious approach, as reflected in the statement *“we choose to be very conservative in our methods”* (R3, 45). However, respondents also indicate that such caution may limit the practical and strategic benefits of AI. As R4 notes, *“some companies are so careful, that they avoid things too far, even though it could actually hurt them to not use it”* (R4, 149), while R6 adds that *“the risks are that you become too stuck in old patterns”* (R6, 209).

Rather than presenting regulation as straightforward, respondents describe a need to interpret where the boundary lies between ethical, responsible, and acceptable AI use. This is reflected in R3's question, *“where's that line of being ethical, being responsible?”* (R3, 79).

5.4 Mitigation strategies

All respondents describe human oversight as a central strategy for mitigating ethical risks in AI-supported recruitment (R1; R2; R3; R4; R5; R6). One respondent explains that *“all AI-generated material is always reviewed and adjusted by me before it is used [...] all decisions regarding candidates are based on my own judgment”* (R5, 159), while another explains that *“I would not trust AI to select my candidates”* (R6, 193).

Respondents also highlight education and continuous learning as important strategies for managing AI-related challenges. One respondent describes that *“We are constantly educating ourselves. However, in order not to be afraid of new tools, we could probably educate ourselves even more”* (R1, 219), while another states that *“there needs to be more education there”* (R4, 143).

What is also emphasized by several respondents, is the strategy of evaluating the systems. This includes maintaining active human involvement, and critically reviewing AI-generated recommendations. As R5 explains, organizations should *“always critically review the content”* (R5, 165). Similarly, one respondent explains that *“we use it within the extent or scope that we have”* (R6, 215), indicating a cautious and controlled approach to AI implementation in recruitment practices. R4 also contributes by talking about the importance to regularly *“evaluate their existing technologies, as well as any new technologies”* (R4, 121), suggesting that ongoing assessment is viewed as necessary to ensure responsible AI use.

Respondents describe practical strategies for reducing bias in AI-supported recruitment. One respondent explains that *“one of the easiest things you can do is make sure that a solution doesn't look at someone's name, even their specific address”* (R4, 108), suggesting that limiting access to personal or demographic information is perceived as a way to reduce the influence of bias in candidate evaluations. The same respondent suggests that *“it's also helpful*

if you have labeled data at some point, so you can go through an exercise where you have humans that are actually evaluating matches between resumes and job descriptions and then you have AI to do the same thing. And then if you have demographic data, you are able to do a comparison between the two, and see the difference between how humans actually evaluate people versus AI.” (R4, 110). The respondent proposes that *"the questions themselves should also be looked at and designed in a way that they're fair and inclusive, and themselves not biased"* (R4, 118).

6 Discussion

6.1 Ethical opportunities

6.1.1 *Efficiency and process support*

The findings indicate that AI creates ethical opportunities through more timely and continuous communication with candidates. This supports previous literature suggesting that AI-supported communication improves candidate engagement, and reduces “time-to-hire” (Muhammad, 2025; Varshney, 2025). From an ethical perspective, timely feedback contributes to fairness and transparency. The findings suggest that AI can create more standardized communication, working independently from recruiters’ workload or time constraints. This may improve candidates’ perception of fairness, since applicants are more likely to receive equal treatment throughout the process, regardless of application timing or volumes.

Several respondents particularly emphasized the value of AI in administrative and repetitive tasks, where automation was perceived as streamlining workflows and reducing manual effort. The findings also show that efficiency influences how decision-making is distributed. Using AI for administrative tasks shifts the recruiters’ focus away from operational tasks, towards evaluative and strategic decision-making. This aligns with Dutta and Naveen’s (2025) concept of “conjoined agency”, where AI and human recruiters work together in the process, rather than replacing one another entirely. Respondents perceive the ethical value of AI as dependent on how it is integrated into recruitment practices, particularly when functioning as decision support, rather than decision replacement.

The findings indicate that AI has the potential to improve efficiency when handling large applicant volumes. R4 describes how recruiters face practical limitations due to time pressure and workload, which may influence how candidates are reviewed. Similarly, R3 describes how applications may be reviewed based on when they applied. This can be connected to Fabris et al. (2026), who found that these constraints may contribute to position bias, where recruiters primarily focus on candidates ranked at the top of a list. In this context, AI is perceived not only as an operational support tool, but also as a potential ethical opportunity. AI may increase the likelihood that more candidates are at least considered within the recruitment process. This supports previous literature arguing that AI enhances recruitment efficiency by processing applications at a scale and speed that exceeds human capability (Varshney, 2025). This suggests that AI may improve fairness not because it is unbiased, but because human recruitment processes are constrained by time, workload, and limited attention.

6.1.2 Reduction of human bias

AI is often described by the respondents as a supporting tool, rather than a replacement for human decision-making. R4 even describes AI as helping them become more aware of their own bias, functioning as a “*second opinion*” and “*watchdog*”. This suggests that AI may contribute to reducing human bias by encouraging more reflection and consistent decision-making processes. It provides a more nuanced perspective to previous research, by indicating that respondents do not necessarily perceive fairness as achieved through removing humans from the process, but rather through balancing human and algorithmic evaluations. In this sense, the findings support Teo’s (2025) argument that fairness in AI systems should be understood as a trade-off between competing objectives, rather than a fixed or fully achievable condition.

The findings also suggest that respondents perceive AI as valuable because it can reduce the influence of interpersonal similarity bias. R2’s reflection that AI may reduce the risk of unconsciously selecting candidates based on similarities, highlights how respondents connect fairness to equal treatment, and standardized evaluation processes. Although R2 does not specifically refer to race, gender, or sexual orientation, the finding can still be connected to previous research showing that human recruitment decisions are often influenced by unconscious preferences, and social perceptions (Pager et al., 2009; Everly et al., 2016). The findings therefore do not necessarily suggest that respondents engage in direct discrimination, but rather recognize the broader risk that subjective human preferences may influence recruitment decisions.

An interesting part of the finding is that respondents acknowledge the limitations of human decision making, while at the same time emphasizing the importance of human involvement. While respondents believe AI can contribute to more objective evaluations, they do not express full trust in automated systems independently making recruitment decisions. This reflects the tension identified in previous research, between algorithmic appreciation and algorithmic aversion (Dietvorst, 2016; Logg et al., 2019). At the same time, some respondents view AI as having future potential to become more consistent, and less biased, than human decision-makers. For example, R4 suggests that AI systems may eventually outperform humans in reducing subjective bias. This suggests that respondents view AI not only as a source of efficiency, but also as a potential tool for improving consistency, and reducing the influence of subjective human judgments in recruitment processes.

On one hand, respondents appreciate AI’s consistency and analytical capabilities. On the other hand, they still perceive human judgment as necessary for contextual interpretation and final decision-making. The findings therefore suggest that respondents position AI as ethically valuable when it complements, rather than replaces, human recruiters. Overall, the findings suggest that respondents perceive AI as ethically valuable. Not because it completely removes bias, but because it may help reduce the influence of subjective human judgments through greater consistency, reflection, and structured evaluation processes.

6.2 Ethical risks

6.2.1 *Introduction to algorithmic bias*

In the findings, algorithmic bias appeared as one of the main risks associated with AI-supported recruitment. Some view AI as being more objective, while others expressed concern that it repeats historically biased patterns.

Since these systems are trained on historical organizational data, respondents fear that existing inequalities and discriminatory patterns may become embedded within algorithmic decision-making processes. This concern reflects previous referenced research, suggesting that algorithmic systems are not inherently objective, but instead shaped by the data and assumptions underlying their design (Köchling & Wehner, 2020). As discussed in the literature review, historical recruitment data may already contain patterns of discrimination related to factors such as race, gender, or socioeconomic background. Consequently, AI systems risk reproducing these inequalities by identifying and favoring characteristics associated with previously successful candidates. This also aligns with Teo's (2025) argument that fairness in AI systems is complex, since efforts to improve efficiency and accuracy may sometimes conflict with goals related to diversity and equal treatment.

Respondents additionally expressed concern that seemingly neutral information, such as names or addresses, may function as proxy variables in AI-supported recruitment processes. R4 highlighted that addresses may indirectly reflect factors such as socioeconomic status, race, or ethnicity, while R3 raised concerns regarding how names may affect evaluation outcomes. These findings indicate that discriminatory patterns may remain embedded within algorithmic systems, even when sensitive variables are formally excluded. These findings correspond with Fabris et al. (2026), who found that non-European candidates may be disadvantaged through indirect signals such as foreign-sounding names, or migration-related indicators, despite having equal qualifications. This suggests that seemingly neutral information may still influence recruitment outcomes, and contribute to discriminatory patterns within AI-supported recruitment.

6.2.2 *Lack of transparency*

Lack of transparency and explainability emerged as another ethical risk in the findings. Several respondents described uncertainty regarding how AI-generated outputs are produced. This aligns with the previous literature on the black box problem, defined as the inability to understand or interpret how the algorithm's logic works when reaching specific decisions (Möller & Petrovic, 2025). This is similar to their study, where they found that recruiters often lacked insight into how AI systems generated outputs and struggled to understand the reasoning behind them. The authors describe how these systems were perceived as opaque, limiting recruiters' ability to interpret or justify algorithmic outcomes.

At the same time, respondents continue interacting with and using these systems, despite lacking full insight into how they operate. This creates an interesting contradiction within the findings, where respondents simultaneously express limited understanding and skepticism toward AI-generated outputs, while still relying on them. This may indicate that practical usefulness and efficiency sometimes outweigh the need for full transparency in organizational recruitment settings.

The findings also suggest that limited transparency may influence how users emotionally and cognitively respond to AI-supported decision-making. This aligns with Tandon et al. (2025) who found that professionals within algorithmic HRM often experience a “fear of the unknown”, due to limited understanding of algorithmic processes. As a result, perceptions of AI-supported decision-making may become ambivalent, where skepticism and appreciation coexist depending on the level of transparency and user understanding.

6.2.3 Potential loss of human oversight and accountability

Regarding accountability, the findings indicate that respondents strongly emphasize the importance of maintaining human involvement in recruitment decisions, particularly by limiting the use of AI in screening and candidate selection. Several respondents describe AI as a support tool rather than a replacement for human judgment. This aligns with previous literature emphasizing that the ethical responsibility for AI-supported decisions ultimately remains with humans (Yanamala, 2023). The findings therefore suggest that human oversight is perceived as essential for ensuring accountability, critically evaluating algorithmic outputs, and maintaining responsibility for recruitment decisions.

Rather than fully trusting AI systems, several respondents described situations where they actively challenged or rejected AI-generated recommendations. They believed the systems could not fully capture candidate qualities or contextual factors. This can be connected to algorithmic aversion, where individuals prefer their own judgment over algorithmic recommendations, despite recognizing the efficiency benefits of AI systems (Dietvorst, 2016). The findings therefore suggest that respondents do not passively accept AI-generated outputs, but actively evaluate and interpret them in relation to their own professional judgment. This further aligns with Karami and Ghasemi’s (2026) argument that algorithmic outputs create the most value when they are interpreted and contextualized by humans. The findings could also reflect Laukkarinen’s (2025) concept of “AI crafting,” where users adapt and reinterpret algorithmic recommendations in order to maintain control over decision-making processes.

However, the findings show an awareness of the risk that human oversight may become superficial or weakened in practice, even though humans try to maintain human accountability. This is mainly discussed by R4, who also describes how time pressures and high workloads may increase the risk of relying on AI-generated outputs. One could argue that this supports Leicht-Deobald et al.’s (2019) discussion of automation bias, meaning

recruiters may only look at the top of the list. Even though all respondents emphasized the importance of keeping humans in the loop, this still suggests that human oversight may become more superficial in practice. This creates a tension between formal human responsibility and meaningful human oversight, meaning it is diffused what human oversight actually qualifies as.

6.2.4 Legal and ethical risks related to data usage

When discussing legal and ethical risks, the findings highlight concerns regarding how candidate information is handled within AI-supported recruitment processes. Respondents describe having only a general understanding of relevant regulations, as well as limited insight into how candidate information is managed internally. This suggests that data governance practices may vary within organizations. It also highlights that uncertainty does not only relate to the AI systems themselves, but also to the organizational practices surrounding data usage. Previous research supports these findings, arguing that ethical challenges in AI are often socio-technical rather than purely technical (Mbiazi et al., 2025)

The findings also indicate that candidates may not fully understand how AI is used during recruitment processes. Respondents describe AI as operating largely “in the background,” making it difficult for candidates to know how their personal data is interpreted or used. This leads to ethical concerns related to informed consent and transparency, in relation to GDPR Article 22. The article addresses automated decision-making with significant effects on individuals. Similarly, Goodman and Flaxman (2017) emphasize the importance of a “right to explanation” in AI-supported decision-making, highlighting the need for individuals to understand how decisions affecting them are made.

It is clear from the respondents that they use third-party tools for AI. However, they express an uncertainty regarding how personal data is managed once entered into external platforms, particularly when organizations lack full control or follow-up mechanisms. This suggests that accountability for data handling may become more complex when AI-supported recruitment depends on external systems and vendors. Even when organizations attempt to follow legal and ethical guidelines, respondents appear uncertain regarding how data is processed, stored, or potentially reused outside the organization’s direct oversight.

6.2.5 Assessment limitations and misinterpretation

The findings show that respondents experience the AI-systems as limited, in terms of accurately interpreting and assessing candidate information. While these issues are primarily technical, respondents highlight that they may lead to unfair or inconsistent outcomes, and are therefore associated with ethical risks. Several respondents described how variations in job

titles, wording, and CV formatting may affect the outputs. This connects closely to the literature on algorithmic fairness, which argues that fairness in AI recruitment is not only about avoiding direct discrimination. It is also about whether systems evaluate candidates accurately and reasonably in practice (Teo, 2025). The findings align with Möller and Petrovic's (2025) study, where the respondents also highlighted that factors such as resume wording and layout can influence algorithmic outcomes. Their research showed that candidates with relevant competencies are not always ranked highly because systems may prioritize formatting and presentation rather than actual competence.

The findings also challenge the assumption that AI automatically improves accuracy and objectivity in recruitment. Previous literature highlights that algorithmic decision-making can process large volumes of data efficiently and produce more consistent outputs than humans (Yanamala, 2021). Similarly, AI recruitment systems are described as tools that improve matching between candidates and job requirements through standardized evaluation criteria (Hawrysz, 2025). However, the empirical findings show that consistency does not necessarily mean correctness. For example, R2 describes how an AI-generated summary incorrectly stated that a candidate lacked Python experience, despite having worked with it. This illustrates how AI systems may produce inaccurate outputs that risk excluding qualified candidates. The findings therefore demonstrate a tension between efficiency and fairness, where faster and more standardized recruitment processes may simultaneously increase the risk of incorrect assessments. However, the findings indicate that HR professionals address these ethical concerns through human oversight and critical evaluation of AI outputs, which was further discussed in the section “6.2.3 *Potential loss of human oversight and accountability*”.

6.3 Ethical ambiguities

6.3.1 *Effect on workforce diversity*

As discussed previously, respondents describe AI as having the potential to reduce certain forms of human subjectivity, through more standardized evaluations and structured recruitment processes. At the same time, both previous research and the empirical findings highlight that AI systems may reproduce historical patterns and existing inequalities when trained on past recruitment data. This creates ambiguity regarding whether AI ultimately contributes to, or limits, workforce diversity.

Several respondents express concern that AI systems may prioritize profiles similar to those historically considered successful. This can potentially reinforce homogeneous recruitment outcomes. For example, respondents discuss the risk that organizations may continue selecting similar candidate profiles, such as “50-year-old white men”, because these profiles reflect previous hiring patterns. However, the findings also suggest that respondents do not necessarily perceive human decision-making as inherently less biased. Instead, some

respondents describe human bias as unavoidable and embedded within recruitment processes, regardless of whether AI is used or not.

This creates an important tension regarding how fairness and diversity should be understood in recruitment. On one hand, respondents are cautious about allowing AI systems to influence candidate selection due to concerns about algorithmic bias and discrimination. On the other hand, some respondents, particularly R4, argue that AI may eventually become less biased and more consistent than human recruiters.

This suggests that respondents do not necessarily view the ethical challenge as a question of choosing between biased humans or biased AI, but rather as a question of how different forms of bias can be balanced and minimized through the cooperation of algorithmic systems and human judgment.

6.3.2 *Perceived fairness*

Fairness in AI-supported recruitment does not appear to be influenced solely by issues related to bias and discrimination. It is also shaped by how competence and candidate suitability are interpreted and evaluated throughout the recruitment process. As aforementioned, Möller and Petrovic (2025) also found that AI systems may prioritize resume wording and layout rather than actual competencies.

Several respondents emphasize the importance of human involvement throughout the recruitment process, particularly because candidate evaluation often requires contextual interpretation. As respondents R3 and R4 explain, recruitment decisions are often “fuzzy”, or “flexible”, rather than straightforward. This means that recruiters must interpret experiences, communication styles, and competence, in relation to the specific role and organizational context. The findings point to a perception of fairness as dependent not only on consistency, but also on flexibility and human judgment.

A central issue in the findings is that AI systems tend to rely on predefined language patterns, keywords, and standardized formats when screening candidates. This suggests that fairness may become dependent on a candidate’s ability to express competence in ways that align with the system’s expectations, rather than on competence itself. As a result, qualified candidates may risk being overlooked if their communication style doesn’t match the logic of the system.

The findings therefore suggest that fairness may be influenced by several layers of interpretation throughout the recruitment process: how competence is expressed by candidates, how it is interpreted by AI systems, and how recruiters understand both the candidate and the algorithmic output. This indicates that fairness is not solely determined by the AI system itself, but also by the interaction between candidates, recruiters, and algorithmic evaluations.

Respondents further suggest that experienced recruiters may be better able to critically interpret AI-generated outputs and identify relevant competence beyond predefined patterns. At the same time, some respondents argue that AI may support less experienced recruiters by helping them structure and evaluate candidate information more consistently. This suggests that AI is not necessarily perceived as more or less accurate than human judgment, but rather as a complementary tool that may support recruiters in critically evaluating candidates and reflecting on their own interpretations.

6.3.3 *The interpretation and application of AI regulation*

Article 22 of the GDPR states that individuals should not be subject to decisions based solely on automated processing when such decisions have significant effects. Organizations may therefore formally comply with this requirement by ensuring that a human remains involved in recruitment processes. However, the empirical findings suggest that the distinction between automated decision-making and human judgment is not always clear in practice.

As R4 explains, even an ordered list without explicit scoring still guides attention toward certain candidates, meaning that the system indirectly shapes decision-making. Although recruiters technically remain responsible for the final decision, AI-generated rankings and recommendations may still influence which candidates receive attention and how they are evaluated.

The empirical findings reflect this ambiguity. Respondents describe regulations as broad and open to interpretation, creating what R4 refers to as “*a lot of grey area*”. Rather than understanding regulation as a clear set of rules, respondents describe a continuous need to interpret what constitutes responsible AI use in practice.

In navigating this uncertainty, some organizations adopt highly cautious approaches to AI implementation in recruitment processes. The findings suggest that uncertainty surrounding regulation and ethical boundaries leads organizations to limit or carefully control how AI is used in practice. At the same time, respondents question whether excessive caution may itself create limitations. The findings suggest that ethical and legal uncertainty does not only restrict potentially problematic AI practices, but may also discourage organizations from using AI in ways that could benefit them.

The findings therefore illustrate that responsible AI use is not perceived as a clear boundary between “allowed” and “not allowed” practices. Instead, organizations appear to continuously negotiate how cautious they should be, balancing regulatory compliance against the perceived strategic benefits of AI-supported recruitment. In this sense, ambiguity emerges not only from the regulation itself, but from the difficulty of determining where responsible caution turns into unnecessary limitation.

6.4 Approaches to mitigate ethical risks

The findings indicate that there are several approaches to mitigating ethical risks. One of the main strategies identified is education and competence development. As R1 explains, their organization is continuously educating their employees, and are aware of the importance of further training related to AI use. Increased knowledge may contribute to a deeper understanding of AI systems and their limitations, which can strengthen critical evaluation, reduce overreliance on automated outputs, and support more ethically informed decision-making throughout the recruitment process.

The findings further suggest that education is closely connected to organizations' ability to critically evaluate both existing and new technologies. Respondents describe how increasing regulatory pressure has led organizations to ask more informed, and critical questions, when procuring AI-supported recruitment systems. According to R4, this has created a “good loop,” where stricter legal and ethical expectations force technology providers to improve areas such as bias testing, transparency, and data handling practices. In this sense, regulation appears to not only constrain organizations, but also encourage more reflective and responsible technology evaluation processes.

At the same time, the findings indicate that understanding how AI systems function is not necessarily the same as understanding how to use them effectively, in practice. R4 describes situations where organizations invest heavily in AI-supported recruitment technologies, but fail to fully utilize them. This suggests that ethical challenges may not only emerge from the systems themselves, but also from how organizations integrate them into existing recruitment practices. The findings therefore highlight the importance of change management and organizational adaptation, when implementing AI-supported recruitment systems.

Another central finding is that respondents consistently emphasize the importance of human involvement throughout the recruitment process. However, the findings also indicate that human oversight is not automatically perceived as sufficient in itself. Rather than simply “keeping humans in the loop,” respondents describe the importance of critical engagement with AI-generated outputs. This supports previous research suggesting that ethical AI use depends not only on human presence, but on active and meaningful engagement with algorithmic recommendations rather than passively accepting them (Karami & Ghasemi, 2026).

7 Conclusion

The study aimed to provide a better understanding of how ethical challenges can be managed, and how AI can be used more responsibly. The research question was: *What ethical considerations emerge when using AI in recruitment, and how is it addressed by HR professionals?* Through a qualitative study based on six semi-structured interviews with HR professionals and individuals working with AI-supported recruitment systems, the findings demonstrate that AI in recruitment creates both ethical opportunities and ethical risks, while also giving rise to several ethical ambiguities.

The findings of this study indicate that the use of AI in recruitment gives rise to several ethical considerations related to fairness, bias, transparency, human judgment, and the handling of personal data. AI is perceived as offering opportunities such as increased efficiency, consistency, and support in administrative recruitment tasks. The findings also highlight concerns regarding biased outcomes, limited transparency, legal restrictions, and the potential loss of human understanding in candidate assessments. The findings therefore suggest that the ethical implications of AI in recruitment emerge through the interaction between organizational practices, humans, and AI systems. Regarding fairness, diversity, and regulation, the findings showed that it was not always clear what should be understood as an ethical opportunity or an ethical risk, which contributed to the identification of several ethical ambiguities and uncertainties surrounding the use of AI in recruitment.

The study further reflects the differences in how HR professionals address the ethical risks. The research shows that this is done to varying degrees, by maintaining human oversight, reviewing AI-generated outputs, limiting AI involvement in high-stakes decision-making, and emphasizing the importance of education and competence development. However, the findings also suggest that these risks may be further mitigated through additional education, clearer legal regulations, and enabling human involvement as a meaningful oversight rather than merely a formal presence in the decision-making process.

7.1 Suggestions for Further Research

This study primarily focuses on recruitment contexts involving relatively standardized application materials and roles within more structured industries. Future research could therefore explore how AI-supported recruitment functions in more creative professions, such as design or other portfolio-based fields, where competence may be expressed in less standardized ways and may be more difficult for AI systems to evaluate through resumes alone.

Future research could also further investigate the relationship between algorithmic bias and human bias in recruitment processes. Several respondents discussed AI as potentially both reducing and reproducing bias, suggesting a need for comparative studies examining how

AI-supported evaluations differ from purely human decision-making. Such research could contribute to a deeper understanding of whether AI systems ultimately reinforce, reduce, or transform existing recruitment biases.

In addition, this study mainly focuses on perceptions and interpretations with AI systems. Future research could therefore examine the technical functioning of recruitment algorithms, including how systems are trained, how recommendations are generated, and how different design choices influence recruitment outcomes.

Appendix

Appendix A. AI Contribution Statement

Tools: ChatGPT, Claude, Whisper

Degree of use: AI-based tools were used as supportive tools throughout the writing process, primarily to improve language, structure, and readability. All generated suggestions and outputs were critically reviewed, edited, and approved by the authors. The tools were used to varying degrees across different sections of the thesis. More extensive support was used in longer and more text-intensive chapters, such as the literature review, while shorter sections, such as the conclusion or introduction, involved more limited use.

ChatGPT and Claude were mainly used for:

- Improving the language and readability of already written paragraphs
- Receiving suggestions on how to restructure sections for better logical flow
- Paraphrasing repetitive or unclear formulations
- Correcting grammar and sentence structure
- Improving transitions between paragraphs and sections
- Suggesting headings based on already written content

Example on prompts:

- “Rephrase this paragraph to make it more cohesive and less repetitive, while keeping the same content.”
- “Suggest a clearer structure for this section.”
- “Improve the transition between these two paragraphs/sections.”
- “Suggest alternative headings based on the content below.”

Whisper was used to transcribe all interviews conducted in the study. The tool made the transcription process significantly faster and reduced the amount of time spent on manual transcription work. This allowed the authors to spend more time on reviewing the material and conducting the thematic analysis of the empirical data. However, all transcriptions were manually reviewed and corrected by the authors to ensure accuracy before the analysis was carried out.

Appendix B. General Interview Guide

Introduction & Use of AI in Candidate Screening/Recruitment Process

- Could you describe your role and your involvement in the recruitment/hiring process
- How, if at all, is AI used in your organization's recruitment process today?

Probes:

- What kinds of tools do you use?
- In which stage(s) of recruitment are they used?
- How long have you been using them?

Efficiency and Consistency

- How do you experience AI improving the efficiency of your work during the recruitment process?
- How consistent is AI in its evaluation of candidates?
- If you have experience with recruitment before AI was used, how does it differ? (e.g., in terms of efficiency)

Interaction & Decision-Making Between Humans and AI

- How do you typically interact with the AI system's recommendations?
 - How do you work with the system's output during screening?
 - What steps do you take after receiving recommendations?
- Do you use your own judgment during the screening process, or do you rely entirely on the AI's evaluation of the candidate?
 - If yes, how? Can you describe situations where you adjust or override the system's recommendations?

Transparency & Explainability

- How much insight do you feel you have into the AI system's algorithms?
 - Can you explain more about how it works?
- Does the AI tool provide information about why a candidate receives a certain ranking or evaluation?

Accountability & Responsibility

- How do you ensure accountability for how candidates are ranked based on your AI system's assessment?

Fairness & Bias in Screening

- Have you experienced any form of bias during the screening process?

- If yes, can you describe such a situation?
 - If no, do you think it could have happened?
- What types of challenges or uncertainties have you experienced when using AI in screening?
- How do you assess whether the screening process is carried out fairly in practice?

Legal and Regulatory Aspects

- How do you ensure that the use of AI complies with regulations such as GDPR or the AI Act?

References

- Agbasiere, C. L. & Nze-Igwe, G. R. (2025). Algorithmic Fairness in Recruitment: Designing AI-Powered Hiring Tools to Identify and Reduce Biases in Candidate Selection, *Path of Science*, vol. 11, no. 4, pp.5001-5021, <https://doi.org/10.22178/pos.116-10>
- Alexander III, L., Song, Q. C., Hickman, L. & Shin, R. (2024). Sourcing Algorithms: Rethinking Fairness in Hiring in the Era of Algorithmic Recruitment, *International Journal of Selection and Assessment*, vol. 33, pp.1-18, <https://doi.org/10.1111/ijsa.12499>
- Björklund, M. & Paulsson, U. (2012). Seminarieboken: Att Skriva, Presentera Och Opponera, 2. uppl., Lund: Studentlitteratur
- Black, J. S. & Van Esch, P. (2020). AI-Enabled Recruiting: What Is It and How Should a Manager Use It?, *Business Horizons*, vol. 63, no. 2, pp.215–226, <https://doi.org/10.1016/j.bushor.2019.12.001>
- Braun, V. & Clarke, V. (2006). Using Thematic Analysis in Psychology, *Qualitative Research in Psychology*, vol. 3, pp.77–101, <https://doi.org/10.1191/1478088706qp063oa>
- Bryman, A. (2018). Samhällsvetenskapliga Metoder, Tredje upplagan., Stockholm: Liber
- Burrell, J. (2016). How the Machine ‘Thinks’: Understanding Opacity in Machine Learning Algorithms, *Big Data & Society*, vol. 3, no. 1, pp.1-12, <https://doi.org/10.1177/205395171562251>
- Cambridge Dictionary. (n.d.a). RACIAL BIAS, Cambridge Dictionary, Available Online: <https://dictionary.cambridge.org/dictionary/english/racial-bias> [Accessed 12 May 2026]
- Cambridge Dictionary. (n.d.b). GENDER BIAS, Cambridge Dictionary, Available Online: <https://dictionary.cambridge.org/dictionary/english/gender-bias?q=gender+bias+> [Accessed 12 May 2026]
- Cambridge Dictionary. (n.d.c). FAIRNESS, Cambridge Dictionary, Available Online: <https://dictionary.cambridge.org/dictionary/english/fairness> [Accessed 12 May 2026]
- Carlsson, M., Finseraas, H., Midtbøen, A. H. & Rafnsdóttir, G. L. (2021). Gender Bias in Academic Recruitment? Evidence from a Survey Experiment in the Nordic Region, *European Sociological Review*, vol. 37, no. 3, pp.399–410, <https://doi.org/10.1093/esr/jcaa050>
- Creswell, J. W. (2009). Research Design: Qualitative, Quantitative, and Mixed Methods Approaches, 3rd edn., Thousand Oaks, Calif: Sage Publications
- Dietvorst, B. J. (2016). Algorithm Aversion, Ph.D., ProQuest Dissertations and Theses, Available Online: <https://www.proquest.com/docview/1807641102/abstract/506008D532BC4BF9PQ/1> [Accessed 20 April 2026]

- Dutta, D. & Naveen, P. M. (2025). Transforming Recruitment and Selection Practices in Organizations Through Discriminative and Generative AI Adoption: A Structuration Lens, *Human Resource Management*, vol. 65, no. 1, pp.77–115, <https://doi.org/10.1002/hrm.70018>
- European Parliamentary Research Service (EPRS). (2024). Artificial Intelligence Act. Author: T. Madiega. PE 698.792. Brussels: European Parliament. Available Online: [https://www.europarl.europa.eu/RegData/etudes/BRIE/2021/698792/EPRS_BRI\(2021\)698792_EN.pdf](https://www.europarl.europa.eu/RegData/etudes/BRIE/2021/698792/EPRS_BRI(2021)698792_EN.pdf)
- European Union. (2016). Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data (General Data Protection Regulation). Official Journal of the European Union, L119, pp.1–88.
- Everly, B. A., Unzueta, M. M. & Shih, M. J. (2016). Can Being Gay Provide a Boost in the Hiring Process? Maybe If the Boss Is Female, *Journal of Business and Psychology*, vol. 31, no. 2, pp.293–306, <https://doi.org/10.1007/s10869-015-9412-y>
- Fabris, A., Rus, C., Saldivar, J., Gatzoura, A., Biega, A. J. & Castillo, C. (2026). Does Fair Ranking Lead to Fair Recruitment Outcomes? A Study of Interventions, Interfaces, and Interactions, *Information Processing & Management*, vol. 63, no. 3, pp.1-18, <https://doi.org/10.1016/j.ipm.2025.104506>
- Floridi, L., Cowls, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., Lütge, C., Madelin, R., Pagallo, U., Rossi, F., Schafer, B., Valcke, P. & Vayena, E. (2018). AI4People—An Ethical Framework for a Good AI Society: Opportunities, Risks, Principles, and Recommendations, *Minds and Machines*, vol. 28, pp.689–707, <https://doi.org/10.1007/s11023-018-9482-5>
- Goodman, B. & Flaxman, S. (2017). European Union Regulations on Algorithmic Decision Making and a “Right to Explanation”, *AI Magazine*, vol. 38, no. 3, pp.50–57, <https://doi.org/10.1609/aimag.v38i3.2741>
- Hawrysz, L. (2025). Artificial Intelligence In Candidate Screening: Opportunities and Challenges, *Scientific Papers of Silesian University of Technology*, vol. 2025, no. 228, pp.203–217, <https://doi.org/10.29119/1641-3466.2025.228.11>
- Herek, G. M., Kimmel, D. C., Amaro, H. & Melton, G. B. (1991). Avoiding Heterosexist Bias in Psychological Research, *American Psychologist*, vol. 46, no. 9, pp. 957-963, <http://doi.org/10.1037/0003-066X.46.9.957>
- Hunkenschroer, A. L. & Kriebitz, A. (2022). Is AI Recruiting (Un)Ethical? A Human Rights Perspective on the Use of AI for Hiring, *AI Ethics*, vol. 3, pp.199–213, <https://doi.org/10.1007/s43681-022-00166-4>
- Hunkenschroer, A. L. & Luetge, C. (2022). Ethics of AI-Enabled Recruiting and Selection: A Review and Research Agenda, *Journal of Business Ethics*, vol. 178, no. 4, pp.977–1007, <https://doi.org/10.1007/s10551-022-05049-6>

- Kalu, F. A. & Bwalya, J. C. (2017). What Makes Qualitative Research Good Research? An Exploratory Analysis of Critical Elements., *International Journal of Social Science Research*, vol. 5, no. 2, pp.43-56, <https://doi.org/10.5296/ijssr.v5i2.10711>
- Karami, M. & Ghasemi, R. (2026). Using AI in HR Decision-Making: What HR Leaders Must Change to Avoid Bias, Blind Automation and Loss of Trust?, *Strategic HR Review*, pp.1–4, <https://doi.org/10.1108/SHR-12-2025-0126>
- Keegan, A. & Meijerink, J. (2025). Algorithmic Management in Organizations? From Edge Case to Center Stage, *Annual Review of Organizational Psychology and Organizational Behavior*, vol. 12, no. pp.395–422, <https://doi.org/10.1146/annurev-orgpsych-110622-070928>
- Köchling, A. & Wehner, M. C. (2020). Discriminated by an Algorithm: A Systematic Review of Discrimination and Fairness by Algorithmic Decision-Making in the Context of HR Recruitment and HR Development, *Business Research*, vol. 13, no. 3, pp.795–848, <https://doi.org/10.1007/s40685-020-00134-w>
- Koivunen, S., Ala-Luopa, S., Olsson, T. & Haapakorpi, A. (2022). The March of Chatbots into Recruitment: Recruiters' Experiences, Expectations, and Design Opportunities, *Computer Supported Cooperative Work (CSCW)*, vol. 31, pp.487–516 <https://doi.org/10.1007/s10606-022-09429-4>
- Krauss, P. (2024). What Is Artificial Intelligence?, in P. Krauss (ed.), *Artificial Intelligence and Brain Research: Neural Networks, Deep Learning and the Future of Cognition*, [e-book] Berlin, Heidelberg: Springer, pp.107–112, Available Online: https://doi.org/10.1007/978-3-662-68980-6_11 [Accessed 8 April 2026]
- Kumares, N., Angelin, R. M., M., D., Pramila, S., Nagrani, K., & Gomathi. (2025). AI in Recruitment: Challenges of Transparency, Accountability and Fairness, *Advances in Consumer Research*, vol. 2, no. 6, pp.2659–2665
- Laukkanen, M. (2025). Working With and Around Artificial Intelligence: AI Crafting and Human-AI Collaboration in Recruitment, *International Journal of Human-Computer Interaction*, vol. 41, no. 21, pp.13804–13816, <https://doi.org/10.1080/10447318.2025.2476715>
- Leicht-Deobald, U., Busch, T., Schank, C., Weibel, A., Schafheitle, S., Wildhaber, I. & Kasper, G. (2019). The Challenges of Algorithm-Based HR Decision-Making for Personal Integrity, *Journal of Business Ethics*, vol. 160, pp.377–392, <https://doi.org/10.1007/s10551-019-04204-w>
- Logg, J. M., Minson, J. A. & Moore, D. A. (2019). Algorithm Appreciation: People Prefer Algorithmic to Human Judgment, *Organizational Behavior and Human Decision Processes*, vol. 151, pp.90–103, <https://doi.org/10.1016/j.obhdp.2018.12.005>
- Mbiazi, D., Bhange, M., Babaei, M., Sheth, I., Kenfack, P. & Kahou, S. E. (2025). Survey on AI Ethics: A Socio-Technical Perspective, *Computational Intelligence*, vol. 41, no. 6, pp.1-24, <https://doi.org/10.1111/coin.70149>

- Möller, N. A. & Petrovic, D. (2025). Artificiell intelligens i rekrytering: En kvalitativ studie om rättvisa, träffsäkerhet och transparens i AI-baserade urvalssystem, Bachelor thesis, Faculty of Technology and Society, Malmö University, <http://mau.diva-portal.org/smash/record.jsf?pid=diva2:1993270> [Accessed 31 March 2026]
- Muhammad, S. (2025). Enhancing Recruiter Outreach: Predicting Job Ad Quality with Advanced Machine Learning and NLP Models, *Artificial Intelligence and Applications*, vol. 3, no. 4, pp.443–458, <https://doi.org/10.47852/bonviewAIA52024083>
- Oates, B. J., Griffiths, M. & McLean, R. (2022). Researching Information Systems and Computing, 2nd edn., Los Angeles: SAGE
- Orbán, F. & Stefkovics, Á. (2025). Trust in Artificial Intelligence: A Survey Experiment to Assess Trust in Algorithmic Decision-Making, *AI & SOCIETY*, vol. 40, no. 6, pp.4955–4969, <https://doi.org/10.1007/s00146-025-02237-6>
- Pager, D., Western, B. & Bonikowski, B. (2009). Discrimination in a Low-Wage Labor Market: A Field Experiment, *American sociological review*, vol. 74, no. 5, pp.777–799, <https://doi.org/10.1177/000312240907400505>
- Raji, I. D., Smart, A., White, R. N., Mitchell, M., Gebru, T., Hutchinson, B., Smith-Loud, J., Theron, D. & Barnes, P. (2020). Closing the AI Accountability Gap: Defining an End-to-End Framework for Internal Algorithmic Auditing, preprint, pp.33-44, <http://arxiv.org/abs/2001.00973>
- Sony, M. M. A. A. M., Amin, M. B., Ashraf, A., Islam, K. M. A., Debnath, N. C. & Debnath, G. C. (2025). Bias in AI-Driven HRM Systems: Investigating Discrimination Risks Embedded in AI Recruitment Tools and HR Analytics, *Social Sciences & Humanities Open*, vol. 12, <https://doi.org/10.1016/j.ssaho.2025.102082>
- Tandon, A., Dhir, A., Malik, A., Budhwar, P. & Kaur, P. (2025). Exploring the Duality of Perceptions: Insights into Uncertainties, Aversion and Appreciation Towards Algorithmic HRM, *Human Resource Management*, vol. 64, no. 2, pp.583–616, <https://doi.org/10.1002/hrm.22263>
- Teo, S. A. (2025). Polycentrism, Not Polemics? Squaring the Circle of Non-Discrimination Law, Accuracy Metrics and Public/Private Interests When Addressing AI Bias, *Frontiers in Political Science*, vol. 7, <https://doi.org/10.3389/fpos.2025.1645160>
- Teo, S. A., Porsdam Mann, S., & Jurcys, P. (2025). The Ethical and Legal Complexities of Regulating Companion AI Chatbots, forthcoming in *Law, Innovation and Technology Journal*, vol. 19, no. 1, SSRN, <https://doi.org/10.2139/ssrn.5778243>
- Varshney, V. (2025). AI Driven Talent Matching Enhancing Recruitment Efficiency And Accuracy, Digital Repository of Theses, [e-journal], Available Online: <https://repository.learn-portal.org/index.php/rps/article/view/1113> [Accessed 15 April 2026]
- Veale, M. & Zuiderveen Borgesius, F. (2021). Demystifying the Draft EU Artificial Intelligence Act, preprint, pp.97–112, <https://doi.org/10.31235/osf.io/38p5f>

- Wachter, S., Mittelstadt, B. & Floridi, L. (2016). Why a Right to Explanation of Automated Decision-Making Does Not Exist in the General Data Protection Regulation, *SSRN Electronic Journal*, vol. 7, no. 2, pp.76–99, <https://doi.org/10.1093/idpl/idx005>
- Wang, X., Zhang, Y. & Zhu, R. (2022). A Brief Review on Algorithmic Fairness, *Management System Engineering*, vol. 1, no. 1, <https://doi.org/10.1007/s44176-022-00006-z>
- Wang, Y. (2025). Algorithmic Decision-Making in Organizations: A Systematic Review toward an Integrated Tension Alignment Framework, *Organization Management Journal*, vol. 23, no. 1, pp.115–131, <https://doi.org/10.1108/OMJ-11-2024-2342>
- Wang, Y., Chen, H., Wei, X., Chang, C. & Zuo, X. (2025). Trusting the Machine: Exploring Participant Perceptions of AI-Driven Summaries in Virtual Focus Groups with and without Human Oversight, *Computers in Human Behavior: Artificial Humans*, vol. 6, <https://doi.org/10.1016/j.chbah.2025.100198>
- Xiong, Y., Kim J. K (2025). Who wants to be hired by AI? How message frames and AI transparency impact individuals' attitudes and behaviors toward companies using AI in hiring, *Computers in Human Behavior: Artificial Humans*, vol. 3, <https://doi.org/10.1016/j.chbah.2025.100120>
- Yanamala, K. K. R. (2021). Integration of AI with Traditional Recruitment Methods, *Journal of Advanced Computing Systems*, vol. 1, no. 1, pp.1–7, <https://doi.org/10.69987/JACS.2021.10101>
- Yanamala, K. K. R. (2023). Transparency, Privacy, and Accountability in AI-Enhanced HR Processes, *Journal of Advanced Computing Systems*, vol. 3, no. 3, pp.10–18, <https://doi.org/10.69987/JACS.2023.30302>
- Yin, R. K. (2016). *Qualitative Research from Start to Finish*, 2nd edn, New York London: The Guilford Press
- Zhang, G., Pan, L., Tang, F. & Yao, F. (2025). Explainable Artificial Intelligence in the Talent Recruitment Process-a Literature Review, *Cogent Business & Management*, vol. 12, no. 1, <https://doi.org/10.1080/23311975.2025.2570881>