

Problematizing the Role of Artificial Intelligence in Hiring and Organisational Inequalities: A Multidisciplinary Review

Forthcoming in *Human Relations*

November 11, 2025

Karen D. Hughes

Department of Sociology
Department of Strategy, Management and Organization
University of Alberta
Visiting Scholar, Babson College
karen.hughes@ualberta.ca

Alla Konnikov

Department of Social Sciences
Concordia University of Edmonton
alla.konnikov@concordia.ab.ca

Nicole Denier

Department of Sociology
Colby College
nicole.denier@colby.edu

Yang Hu

UCL Social Research Institute
University College London
yang.hu@ucl.ac.uk

This is the author accepted manuscript for publication. The final version will be available in *Human Relations*. The authors wish to thank the handling editor and anonymous reviewers for their helpful feedback throughout the review process.

Abstract

What are the implications of the growing use of artificial intelligence (AI) in recruitment and hiring for organisational inequalities? While advocates suggest that AI is a groundbreaking tool that can enhance hiring precision, efficiency, diversity, and fit, critics raise serious concerns around bias, fairness, and privacy. This review article critically advances this debate by drawing on diverse scholarship across computing and data sciences; human resource, management, and organisation studies; social sciences; and legal studies. Using a hybrid review approach that combines scoping and problematising review methods, we examine the implications of algorithmic hiring for organisational inequalities. Our review identifies a multidisciplinary discussion marked by asymmetries in how key concerns are conceptualised; a clear and heightened potential for AI to conceal inequalities in hiring processes; and contestation over the regulation of algorithmic hiring. Building on Acker's (2006) framework of 'inequality regimes', we propose the concept of *algorithmically-mediated inequality regimes* to highlight AI's capacity for concealing and reproducing inequalities in hiring through enhanced *algorithmic invisibility* and the growing *legitimacy* of AI solutions. We propose an agenda for future research, policy, and practice, emphasising the need for an interdisciplinary 'chain of knowledge' and a multi-stakeholder 'chain of responsibility' in AI application and regulation.

Introduction

Hiring and recruitment practices are being dramatically reshaped by artificial intelligence (AI), sparking important debate over the implications for individual workers and organisational inequalities (Bornstein, 2018; Nawaz, 2019; Upadhyay and Khandelwal, 2018). Advocates view AI as an enticing ‘technosolution’ to perennial challenges in hiring, concerning accuracy, fit, efficiency, and authenticity (Roemmich et al., 2023). Critics raise serious concerns about issues of bias, fairness, and privacy (Aizenberg and Van Den Hoven, 2020; Burrell and Fourcade, 2021; Weiskopf and Hansen, 2023). While hiring processes have long relied on technologies—such as databases, resume screening software and cybervetting (Berkelaar, 2017; Friedman and McCarthy, 2020)—it is clear that AI is ushering in a fundamentally new era. Today’s complex algorithmic ecosystems operate at unprecedented speed and scale, optimising job postings, screening resumes, and analysing candidates’ skills, body language, and speech to predict ‘ideal’ matches (Ajunwa and Green, 2019; Ajunwa, 2021b; Bornstein, 2018; Kelan, 2023; Manroop et al., 2024; Weiskopf and Hansen, 2023). Yet, while AI promises benefits—streamlining hiring processes for employers, and simplifying job searches for applicants—it also carries significant risks, potentially amplifying bias and obscuring the mechanisms that sustain inequality.

Hiring is a critical and theoretically rich site for considering the implications of AI. It is widely acknowledged as a foundational gatekeeping mechanism within market-based capitalist societies that precedes and conditions all subsequent organisational processes (Acker, 2006; Kim, 2020). Unlike other HR functions, impacting those already employed in organisations, hiring processes determine who is included and who is excluded, thus structuring economic opportunities and ‘life chances’ (Weber, 1978). Inequality produced in hiring processes has unique stakes, with the capacity to generate cascading cumulative effects on career trajectories,

income, and social mobility (DiPrete & Eirich, 2006; Rivera, 2020). Individuals *not* hired cannot be evaluated or promoted, and are thus excluded from the opportunities to perform and compete for rewards within organisations (Bills et al, 2017). As AI is rapidly integrated into hiring processes (SHRM, 2022), shaping access to rewards and opportunities, it is timely to examine whether and how AI is (re)shaping inequality.

Our work offers a multidisciplinary review of scholarship and debates currently grappling with the implications of AI for hiring inequalities. Ongoing debates are occurring across varied disciplines—including computing science, human resource management, the social sciences, and law—highlighting that AI in hiring is an inherently multidisciplinary phenomenon. Yet, to the best of our knowledge, no review has explored this issue from a multidisciplinary perspective (e.g. Bankins et al., 2024; Basu et al., 2023; Burrell & Fourcade, 2021; Chen, 2023b; Kellogg et al., 2020). Embracing a multidisciplinary view is particularly important for understanding the role of AI in hiring, where technical, social, organisational, and legal considerations intersect to shape opportunities. As our review demonstrates, despite shared concerns, scholars often work in disciplinary silos, guided by distinct foci and assumptions, with limited exchange. This lack of cross-fertilisation hinders a comprehensive understanding of how AI contributes to the (re)production of inequalities.

Theoretically, we employ Acker’s (2006) conceptualisation of ‘inequality regimes’ to underscore the structural pervasiveness of organisational inequality and to problematise debates occurring across and within disciplines. Although Acker’s writing predates the rise of AI, her perspective is highly relevant to algorithmic hiring, where inequalities can be embedded within data, design, and the broader ecosystem. This approach underscores how multiple inequality-

producing mechanisms can pervade routine organisational processes, remaining invisible while appearing legitimate and difficult to challenge.

Methodologically, we undertake a hybrid approach to reviewing and bringing this multidisciplinary scholarship together. First, we assess the breadth of the emerging debates using a scoping methodology (Arksey and O'Malley, 2005) to identify relevant scholarship in a rigorous manner. Specifically, we ask: *What key questions and concerns are being explored across disciplines regarding the role of AI in organisational hiring processes?* Second, we assess the depth of the emergent foci, ideas, and debates via a problematising review (Alvesson and Sandberg, 2020) which questions underlying assumptions and frameworks in order to enable the development of new ideas and concepts. Specifically, we ask: *What underlying assumptions and approaches guide each discipline?* In answering these questions, we bring Acker's framework of 'inequality regimes' (2006) into the digital era to problematise current debates and to conceptualise the place and the role of AI in (re)shaping and interacting with organisational inequality regimes, focusing on how its increasing pervasiveness may reinforce, reconfigure, or obscure inequalities.

Our multidisciplinary review makes several contributions to advancing knowledge about hiring and labour market inequalities in the digital age. First, we demonstrate that despite shared concerns about AI's potential to exacerbate organisational inequality, there is a growing dominance of technical knowledge and perspective which risks marginalising critical, alternative perspectives. Second, the current framing of AI as a solution to human bias narrows the focus to individual decision-makers and technical solutions, neglecting the pervasiveness of structural inequalities that are becoming further concealed within what we conceptualise as *algorithmically-mediated inequality regimes*. Third, we note that regulatory responses to AI are

characterised by significant lag and contestation, with outcomes varying across global, national, regional, and local contexts. This complicates the development of coherent strategies to address AI's far-reaching implications for hiring. It is only through a multidisciplinary lens that these patterns, tensions, and blind spots become visible. We situate our analysis within broad structures of power asymmetries, exacerbated by AI and digitisation, recognising that organisations do not operate in a vacuum. We conclude by outlining an agenda for future theorising, research, and practice. We emphasise the need for interdisciplinary dialogue and multi-stakeholder 'chain of responsibility' to foster a more holistic understanding of AI's transformative effects.

Conceptual Framework: Inequality Regimes and AI

We extend Acker's (2006) framework of inequality regimes to critically examine the implications of AI for hiring inequalities. 'Inequality regimes' are defined as 'interrelated practices, processes, actions, and meanings that result in and maintain inequalities within organisations' (443). Acker (2006) identifies hiring as a pivotal organisational process that structures access to opportunities and reinforces disparities across class, gender, and racial lines. AI is defined as computational systems that simulate human intelligence to perform tasks such as learning, reasoning, and decision-making (Russell and Norvig, 2021). In hiring, AI encompasses technologies including natural language processing, machine learning, and facial recognition, among others, which are used to optimise job postings, analyse resumes, and assess candidates' speech and behaviour.

The use of AI in hiring both echoes and extends Acker's theorisation of inequality regimes. Acker (2006) argues that inequality-producing mechanisms often operate invisibly

within routine processes, seemingly legitimate and difficult to challenge. It is only when they become visible that their legitimacy is questioned. Yet, AI solutions often operate in the opposite direction—first, with almost ubiquitous legitimacy, as evidenced by the growing normalisation of digital trace collection and usage (Elliott, 2019; Zuboff, 2019), and second, with heightened invisibility, where AI functions as a ‘black box’ that renders decision making processes more opaque and difficult to trace (Ajunwa, 2020b; Pasquale, 2015). These changes are situated within broader systemic changes conceptualised by Zuboff (2019) as *surveillance capitalism*—a new economic logic in which digital tracing and data extraction serve to generate profit and intensify power asymmetries between corporations and individuals. As Pasquale (2015) and Zuboff (2019) argue, these effects are not mere byproducts of the technology but are intentional, reproducing and magnifying power asymmetries.

Hiring offers an ultimate example of how these power dynamics operate (Ajunwa and Green, 2019). Despite being a two-sided process, power relations within hiring are highly asymmetric. Employers make high-stakes decisions defining who will be granted access to a job opportunity, a key life-chance structuring experience within market-based, capitalist, economies. While hiring is driven by the search for the ‘ideal worker’, the stated criteria may not be neutral but gendered, racialised and intersectional, and shaped by understandings of class and social location (Acker, 2006). Despite efforts to the contrary, hiring procedures and outcomes often remain enigmatic, leaving job seekers struggling to adapt their profiles and maximise their odds of being seen as an ‘ideal candidate’ (Ajunwa and Green, 2019). The increasing automation and augmentation of hiring processes exacerbates these complexities, posing new concerns (Kellogg et al., 2020; Newman, et al., 2020). Amongst these are the growing legitimacy of complex algorithmic hiring tools, with magnified, often unregulated access to data (Zuboff, 2019); the

increased opaqueness of decision-making (Pasquale, 2015); and growing power asymmetries in hiring (Ajunwa, 2020a). This constrains the agency of job seekers, limiting the choices and ability to understand or shape decisions made about them.

Despite this fundamental transformation of hiring, the role of AI in (re)shaping hiring inequality has not yet been thoroughly examined. Important reviews of organisational inequality, such as Amis et al. (2020), discuss hiring's central role in producing inequality but omit considerations of AI. Our multidisciplinary review aims to clarify the relationship between AI, hiring, and inequality. Extending Acker's framework to the digital era helps to problematise the way technological developments intertwine with inequality regimes.

Methodology

Our positionality grounds how we engage with this review. As a team of four social scientists with distinct backgrounds, and shared expertise in social, organisational, and labour market inequality, we bring three vantage points to this review. First, we view social inequality as a structural rather than individual phenomena. Accordingly, work and organisations are structures in which inequality is produced and reproduced. Second, inspired by constructivist perspectives (Alvesson and Sköldberg, 2018), we view knowledge as socially constructed, not neutral, and shaped by power dynamics. Third, through shared experiences of conducting research on the implications of AI for labour market inequality as a part of a multidisciplinary team (including computing, statistics, and social science scholars), we are practically attuned to how different fields study this issue. This informs our attention to how knowledge is produced, how certain ideas gain traction while others are underexplored, and how disciplinary assumptions, including our own, shape knowledge.

Similarly, we recognise that our interpretive choices are shaped by our epistemological stances and disciplinary assumptions. For instance, perceiving social and organisational inequality as a structural phenomenon makes us more attentive to whether and how inequality is defined and discussed in different bodies of scholarship, and to whether these discussions occur at structural or individual levels. Our approach to knowledge formation also made us attentive to the language and terminology used to discuss inequality, and to how language reflects underlying disciplinary assumptions (e.g., ‘bias’ and ‘debiasing’ signalling a technosolutionist approach). Our experience of working in the multidisciplinary team has taught us to look within and across disciplines, building deep intra-disciplinary knowledge and creating ‘bridges’ through inter-disciplinary perspectives.

Our hybrid methodology, combining scoping and problematising methods, highlights our explorative stance, focusing on *what* knowledge is being created (Arksey & O’Malley, 2005), and our constructivist stance, focusing on *how* knowledge is being created (Alvesson and Sköldbberg, 2018). Ultimately, our aim is not simply to map key debates and participating disciplines, but to critically question how knowledge is being shaped (see Supplemental Materials C for further details). This hybrid methodology is essential for addressing the complexity of AI and hiring research as it evolves rapidly across disciplines. A key strength of a scoping review is its capacity to assess the breadth of a growing body of work that defines the ‘literature’ and to map emerging ideas and questions. The problematising review complements and deepens this by critically engaging with key debates and assumptions. Together, this hybrid approach (see Figure 1) supports our aim of fostering interdisciplinary dialogue on AI’s role in hiring and its implications for organisational and social inequality.

Phase 1: Scoping review

A scoping methodology is ideally suited for mapping emerging, complex, and evolving fields (Mays et al., 2001: 194). It provides flexibility to incorporate work from distinct disciplines and traditions (Daudt et al., 2013) rooted in an explorative, question-driven approach to knowledge creation, rather than a hypothesis-oriented one. While a scoping review methodology does not incorporate explicit quality assessment—a key distinction from the systematic review—it is nonetheless a rigorous, multi-step process aimed at maximising breadth and inclusion of non-mainstream scholarship (Arksey & O'Malley, 2005). The steps are presented on the left-hand side of Figure 1.

In the first step, we worked to identify relevant literature. Guided by our research questions, we used broad criteria to collect relevant items at the intersection of the three domains of focus: 'AI', 'hiring', and 'inequality'. We operationalised this via the combination of the following keywords: 1) *AI and algorithms*; 2) *human resource management, hiring, recruitment, work, employment and organisational processes*; and 3) *(in)equality, fairness, bias, stereotypes, prejudice and discrimination*. To ensure that we captured relevant pieces, we searched multiple databases, starting with Google Scholar, then ProQuest, Sociological Abstracts, Social Sciences Citation Index, Sage Journals, and Social Sciences Research Network, until no new items emerged. Our initial search resulted in 389 publications.

In the second step, we screened publications, using the Rayyan platform to allow for a blind screening process by our four-author team. We coded publications as: 1) 'highly relevant', 2) 'somewhat relevant', and 3) 'least relevant' to our focal questions (see Supplemental Materials C, page 4 for details). Following majority agreement (three or more team members) of

75%, we included 212 ‘highly relevant’ publications (i.e. central focus on and sustained analysis of the intersection of ‘AI’, ‘hiring’, and ‘inequality’).

[Figure 1 here]

Our third step involved a detailed mapping of these 212 publications, noting discipline, key questions, concepts, and insights (see Supplemental Materials A for the coding system). Using the combined information on the journal, the authors’ background and disciplinary affiliation, along with the keywords and the abstract, we identified four broad disciplinary clusters connected to: 1) computing and data sciences (CS); 2) human resource, management, business and organisation studies (HRMOS); 3) social sciences (SS); and 4) legal scholarship (LS). Given the complexity and novelty of emerging knowledge, we do not view these groupings as rigid though we observed distinct characteristics in each cluster. For instance, CS research on AI is vast; we thus limited our focus to publications with a clear focus on the social implications of AI. HRMOS scholarship is very diverse, with applied and conceptual writings. Social sciences (SS) are the broadest, most heterogeneous group, encompassing sociology, psychology, philosophy, and socio-technical orientations. Legal scholarship (LS) offers more focused writing on specific legal questions. In a small number of cases where an item could potentially belong to more than one cluster (i.e. socio-technical perspectives which draw expertise from the CS and SS), we assigned it to the cluster that was most substantively relevant (see Supplemental Materials C for further details).

Phase 2: Problematising review

While a scoping strategy is essential for identifying this emerging multidisciplinary body of work, the value of a problematising approach is that it critically assesses how knowledge is being

constructed, unpacking the underlying assumptions that shape knowledge creation, and offering new insights and questions as a result. It therefore differs from familiar review approaches (e.g. integrative, systematic) that seek to offer a ‘representative’ description of knowledge, aiming instead to ‘re-evaluate existing understandings of phenomena, with a particular view to challenging and reimagining our current ways of thinking about them’ (Alvesson and Sandberg, 2020: 1297). We followed four key principles of the problematising approach: 1) using researcher reflexivity as a resource; 2) recognising that ‘less is more’ by focusing on the most substantively insightful contributions; 3) ‘reading broadly but also selectively’; and 4) working to problematise rather than accumulate knowledge (Alvesson and Sandberg, 2020: 1300) (see Supplemental Materials C for further details).

Reflexivity is a key element that drives problematisation. Reviewers are not perceived as ‘neutral’ but active actors whose ‘intellectual resources’ and ‘paradigms and fashions’ are essential in analytical processes (Alvesson and Sandberg, 2020: 1295, 1297). Given the team-based nature of our project, we refined the problematising review principles to incorporate processes of co-reading and dialogue, moving from individual to more collective forms of knowledge creation (Mauthner and Doucet, 2008: 977). Specifically, we conducted numerous layered co-reading exercises, assigning a primary reader to each discipline to allow the accumulation of expertise in the specific domain, and secondary readers who rotated across disciplines, allowing for the generation of cross-disciplinary ‘bridges’. The depth and reflexivity that this type of co-reading enabled were essential to developing insights within and across disciplines concerning the key questions, concepts, and assumptions in the scholarship. These insights became a foundation for problematising.

Our strategy of assigning a ‘primary reader’ to each disciplinary cluster to accumulate knowledge and a ‘secondary reader’ to rotate across disciplines fostered expertise in Phase 1 that allowed us to refine the initial corpus, and target a more select subset of relevant publications in Phase 2. Following the principle of ‘less is more’, we created a ‘supercorpus’—a selection of 97 readings, from the scoping review (57 items) and additional readings from reading ‘broadly but selectively’—beyond the corpus (40 items). We identify items cited from this final corpus in our bibliography with an asterisk [*]. A summary of the corpus, with cluster, key foci, and key concepts appears in Supplemental Materials B.

To identify scholarship for the ‘supercorpus’ we relied on each reviewer’s accumulated ‘intra-disciplinary’ knowledge, developed through immersive reading and engagement with their focal area. Our team’s collective intra-disciplinary knowledge thus functioned as a form of quality assessment. Our selection criteria was designed to identify publications that: 1) offered the most substantive analysis, 2) sparked key debates, 3) introduced novel concepts, and/or 4) demonstrated high relevance to the intersection of AI, hiring, and inequality. In line with Alvesson and Sandberg (2020: 1297), we gave priority to conceptually generative insights. We also read beyond the initial corpus—a common practice in the problematising method—by hand searching for key classic and contemporary texts to aid our analysis. While in Phase 1, our selection criteria were strongly tied to the intersection of three domains (i.e. AI, hiring, and inequality), in Phase 2, other items in neighbouring or overlapping domains (e.g. digitalisation) were considered to be relevant for a deeper understanding of current debates (further details appear in Supplemental Materials C, page 7).

We then reviewed these 97 items using deep, reflexive, layered reading. This follows Alvesson and Sandberg’s (2020) view of problematising as ‘an “opening up exercise” that

enables researchers to imagine how to rethink existing literature in ways that generate new and “better” ways of thinking about specific phenomena’ (1290). We use Acker’s (2006) classic work on inequality regimes as a valuable framework for rethinking the potential of AI to contribute to organisational inequalities. We also incorporate contemporary scholarship on the broader implications of AI and digitisation, including Elliott’s book (2019) on the growing use and cultural acceptance of AI, Zuboff’s work (2019) on surveillance capitalism, data, privacy, and control, and Pasquale’s (2015) insights on the intentional complexity and opaque nature of AI, all of which set the stage for organisational use of AI.

In the following sections, we discuss questions and concerns examined across various disciplines (RQ1) and uncover underlying assumptions within different disciplines (RQ2). In the discussion section that follows, we problematise and explore novel ways of conceptualising AI’s role in hiring, situating the analysis within Acker’s (2006) framework of inequality regimes. We conclude by considering the implications of our findings and offering an agenda for future research. Table 1 (below) summarises our key findings.

Results: The state of the multidisciplinary field

RQ1. What key questions and concerns are being examined in distinct disciplines about the role of AI in organisational hiring and recruitment processes?

Our review begins by examining key questions posed across disciplinary clusters, recognising the central role questions play in shaping knowledge production (Alvesson and Sandberg, 2000: 1294). While questions about bias, fairness, discrimination, transparency and accountability are shared, disciplines differ in how these issues are defined and the solutions they propose. Computing science (CS) emphasises technical causes and fixes of ‘bias’; human resource and

management scholarship (HRMOS) focuses on trust and improving the hiring process; and the social sciences (SS) and legal (LS) scholarship raise broader social, ethical, regulatory, and political questions. We begin with the CS literature which drives the technological phenomenon we are studying and dominant ways of framing the problem.

[Table 1 here]

Computing and data science (CS): Broadly speaking, CS focuses on technical solutions (Lepri et al., 2018) and, in some cases, the social-technical interface (Amershi et al., 2014). The core interest is in developing AI to enhance organisational efficiency, accuracy, and authenticity in hiring, and measuring and mitigating bias (Glymour and Herington, 2019; Kuhlman et al., 2020). CS examines questions of inequality through the lens of ‘bias’, defined as systematic deviation between expected and actual values of the predicted outputs of an AI application, resulting in unequal (dis)advantage for individuals or groups (Kordzadeh and Ghasemaghaei, 2022). The concept of *demographic parity* (i.e. proportional representation of different groups) is commonly used as ‘ground truth’, with any deviation in AI algorithms or outcomes considered ‘bias’ (De Alford et al., 2020; Geyik et al., 2019).

Specific questions of interest in CS thus focus on technicalities regarding: 1) how to select and debias datasets used for training AI algorithms; 2) how to optimise AI algorithms to reduce bias in AI outputs (e.g. feature removal from models) (Geyik et al., 2019); 3) who should be involved in the debiasing process (e.g. AI developers, clients, end-users) (Amershi et al., 2014; Birhane and Cummins, 2019; Raghavan et al., 2020); 4) why bias mitigation techniques fail (Gonen and Goldberg, 2019); and 5) if it is not possible to eliminate bias, how AI developers can instead focus on ‘procedural justice’, ensuring that AI development and deployment processes are fair, transparent, and accountable (Lepri et al., 2018).

Organisations are a frequent focus of empirical investigations in CS for several reasons. First, they are sites of high-stakes AI deployment, where decisions can significantly impact individuals and risk breaching anti-discrimination laws (Baer, 2019). Second, organisational settings reveal visible forms of bias, such as disparities in workforce composition, wages, and job satisfaction (Lee et al., 2015). Finally, they offer accessible data for AI training and bias analysis. Yet, organisations are often treated in a generic manner, with limited attention to their specific dynamics, structures, or processes.

Human resource, management and organisation studies (HRMOS): Much writing in this stream focuses on the practical application of AI tools in recruitment and selection, and the potential for more ‘efficient’ and ‘effective’ organisational processes as a result (Allal-Chérif et al., 2021; Upadhyay and Khandelwal, 2018). A first set of questions concern whether and how AI solutions, such as video interviews, chatbots, and automated screening tools, can improve hiring processes, making them faster, less costly, more predictable, and less biased (Dattner et al., 2019; Newman et al., 2020). AI is often perceived as a solution to human bias, frequently conceptualised as ‘implicit’, suggesting the unintentional bias generated as a result of the complexity involved in human decision making.

A second set of questions concern how AI can address uncertainty and complexity in organisational decision making (Dwivedi et al., 2019). Debate here focuses on whether AI will replace (*automation*) or complement (*augmentation*) human decision making (Dwivedi et al., 2019; Jarrahi, 2018; Raisch and Krakowski, 2021). There is a strong argument for AI-human collaboration, with AI handling routine tasks such as resume screening, while humans tackle higher-order, tacit knowledge tasks (Bankins et al., 2024; Dwivedi et al., 2019; Jarrahi, 2018;

Shrestha et al., 2019). Huang and Rust (2018) suggest AI handles mechanical and analytical tasks, leaving intuition and empathy to humans. These discussions, however, rarely consider the ethical aspects of such collaboration (Bankins et al., 2024).

Trust is a third concern: that is, whether AI-based decisions are trustworthy and trusted within organisations, by applicants, and by the public, and what conditions can enhance perceptions of fairness. While AI promises to reduce human bias, applicants prefer and trust human ‘two-way’ interactions over machine-based ones (Acikgoz et al., 2020; Newman et al., 2020). HRMOS explores solutions to enhancing trust, emphasising transparency, fairness, and algorithmic oversight (Langenkamp et al., 2020). A key concern is how algorithms trained on non-representative or skewed datasets create inaccuracies and bias (Tambe et al., 2019). Ethical and legal concerns are also raised, including privacy issues when using training data without consent (Dattner et al., 2019) and questions about which parties—developers or employers—bear responsibility for ethical outcomes (Martin, 2019).

Social sciences (SS): Encompassing diverse fields, the SS provide a broader critical perspective on AI and hiring, incorporating varied theories and levels of analysis. One important area of interest is on *affordances*—that is, what AI enables in hiring. For instance, Ajunwa and Greene (2019) describe how automated hiring platforms facilitate the fungibility of workers within and between organisations, capturing and structuring data in ways that aid employer control and standardised management. Another issue concerns the bundling of diverse data (e.g. employer records, social media, personality assessments, behavioural traces) and heightened potential for novel forms of hiring discrimination (Cruz, 2024; Rosenblat et al., 2014; Zuboff, 2019). For instance, Cruz (2024) shows how AI and social media are used to collect and assess non-skill

related information, such as political leanings or the willingness to relocate, highlighting concerns over privacy, bias, and discrimination.

If emerging technologies afford managers and employers enhanced power—through access to more information, sophisticated surveillance, and advanced analytics—what are the implications for workers and job seekers? Researchers are beginning to document how and whether discrimination occurs in algorithmic hiring and evaluation (Ajunwa, 2020b), and how individuals perceive and respond to these technologies (Gelles et al., 2018; Lee, 2018).

Individuals may engage with, evaluate, or resist these systems. Studies raise concerns on how workers perceive and react emotionally to algorithmic hiring and management systems with attention to perception of fairness and trustworthiness (Gelles et al., 2018).

Given these concerns, there is also attention to what can or should be done to regulate new technologies, or use them in ways that promote socially desirable outcomes. Some propose frameworks to mitigate bias through technical design improvements in AI systems (Lin et al., 2020). Others advocate for stakeholder engagement, such as the ‘Design for Values’ approach which promotes designing socio-technical systems that support human rights (Aizenberg and Van Den Hoven, 2020). Unions and worker representatives are also recognised as key actors in protecting or securing rights, given their role within many workplaces (Todoli-Signes, 2019). Broader organising efforts—beyond specific workplaces—are increasingly seen as crucial for shaping the development and use of AI technologies in ways that safeguard worker interests (Crawford et al., 2019). Legal frameworks are another focus in addressing the challenges posed by AI technologies (Rovatsos et al., 2019), particularly in relation to human rights, privacy, and data regulation. These concerns bridge social sciences and legal scholarship, particularly in employment, where non-discrimination is a core legal protection in many capitalist democracies.

Legal scholarship (LS): Central areas of attention include bias, discrimination, transparency, and privacy, with particular attention to the potential harms faced by individuals navigating ‘opportunity markets’ (Ajunwa, 2020a; Kim, 2020). These concerns intersect with broad debates over AI global governance, such as the three competing models discussed by Bradford (2023) in *Digital Empires*: the US market-driven approach, which prioritises corporate interests; the EU rights-driven model, emphasising individual protections; and the Chinese state-driven model, focused on centralised control. Most studies in our review focus on US and EU contexts, reflecting the dominance of a Global North perspective and exclusion of the Global South. In the US, legal attention typically centers on employment law, particularly Title VII (The Civil Rights Act of 1964), which prohibits discrimination on protected grounds, and the lack of proactive measures and limitations of complaint based approaches. In contrast, writing on the EU emphasises its preventative, rights-based approach. For example, the General Data Protection Regulations (GDPR) protects personal data, restricts solely automated decision-making, and mandates human oversight in high stakes areas like hiring (Kaminski, 2018)—principles reinforced and extended in the EU’s 2024 AI Act, which we discuss later.

A recurring question in legal scholarship is whether existing frameworks are a match for rapidly evolving technologies that make it difficult or impossible to detect bias (Ajunwa, 2020a; Hacker, 2018; Kim, 2020). Bent (2020) questions whether laws ‘developed to govern humans’ can ‘translate readily to the government of machine-assisted decisions’ (805). Mann and Matzner (2019) warn that ‘with increased algorithmic complexity, biases will become more sophisticated and difficult to identify, control for, or contest’ (1). Ajunwa (2020a; 2020b) underlines how algorithms can radically amplify bias, extending the traditional scope of harm, and affecting vast numbers of people.

Such concerns are accompanied by specific questions about how established law might address algorithmic bias and discrimination. In the US, scholars debate whether legal concepts of ‘disparate treatment’ and ‘disparate impact’, and complaint-based versus proactive approaches, can inform evolving practice. If biases are unobservable, and unknown, making complaint-based processes impossible, is proactive regulatory intervention at the design stage not essential (Kim, 2020)? Does failure to consider disparate impact constitute evidence of discriminatory intent and should proactive auditing approaches be mandatory (Ajunwa, 2020a)? Does achieving ‘fairness’ in employment decisions by accounting for protected characteristics constitute ‘algorithmic affirmative action’ and, if so, is this legal (Bent, 2020)? And since humans are involved in AI production and use, who is legally accountable for the inequalities that algorithms produce given evolving relationships between employers, developers, vendors, and platforms (Ajunwa, 2020a; 2020b)?

Across jurisdictions, broader field-level questions emerge about stakeholder roles and governance, helping to connect key questions within this review. Do effective responses to algorithmic bias require ‘technical’ fixes (e.g. debiasing algorithms), legal reforms (e.g. improving existing anti-discrimination legislation), human oversight (e.g. third party auditors), or a combination of socio-legal-technical solutions (Bornstein, 2018; Nachbar, 2021)? How should countries balance individual rights, business necessity, and systemic regulation (Kaminski, 2018)? And given the limitations of existing employment and labour law, what other areas of law—for instance, privacy law and consumer protection—might be mobilised into a multi-pronged legal approach (Hacker, 2018; Mann and Matzner, 2019)?

RQ2. What are the underlying assumptions and approaches within different disciplines concerning algorithmic hiring and its relationship to organisational inequalities?

Building on our hybrid review approach, we problematise to think critically about the implicit / explicit assumptions and established conventions in each field. Overall we find that disciplinary assumptions vary notably. CS assumes bias is a technical problem that can be isolated and solved; HRMOS assumes AI can help to improve human decision making; SS assumes AI is entangled with broader systems of inequality; while much LS assumes AI can be governed through some type of legal regulations.

Computing and data science (CS): A distinguishing feature of CS is its limited substantive engagement with organisational contexts and processes. Typically organisations are viewed as data sources and ‘high-stakes’ domains that can elevate the importance of AI technical developments / applications. Despite this narrower lens, CS makes clear assumptions about how AI contributes to both the (re)production and mitigation of organisational inequalities. Specifically, it assumes the two key sources of bias are: 1) data and 2) algorithms (i.e. the methods used to process data).

For the former, issues such as inherent bias in the data used for algorithm training are core to the computing literature—specific aspects include data scarcity or missing data on underrepresented groups (e.g. racial minorities) and imbalanced / biased data representation of different social groups (e.g. the under-representation of racial minorities in training data) (Favaretto et al., 2019; Kuhlman et al., 2020). Such issues lead to further disadvantages for marginalised groups in AI applications such as CV screening and AI-powered job interviews and competency assessments (Geyik et al., 2019; Langenkamp et al., 2020).

Algorithms are often seen as the primary means for rebalancing and rectifying biases in training data. Common techniques include feature removal (e.g. excluding race as a predictor of performance in job applicant assessments), data reweighting (e.g. increasing the weighting of under-represented groups), and imposing external rules for calibrating algorithm outputs (e.g. the four-fifths rule whereby the proportion of candidates from a minority group should not be below 80% of the majority group) (Raghavan and Kim, 2023). However, some CS studies argue that algorithms can only address a limited number of known biases (i.e. characteristics such as gender and race explicitly chosen by users / designers) and cannot ‘remove’ biases embedded in complex latent patterns in training data that are not known or specified as a source of bias (Gonen and Goldberg, 2019). In short, algorithmic solutions cover up rather than remove biases that lead to AI-induced inequalities (Gonen and Goldberg, 2019). Recognising that ‘seeing’ and ‘knowing’ biases is key to identifying and mitigating them, CS literature advocates for more diversity in the AI workforce (Kuhlman et al., 2020).

Human resource and management studies (HRMOS): Though often focused on AI implementation in hiring processes, scholarship in this stream reflects several underlying assumptions. First, it is assumed that organisational processes in general, and hiring processes and decision making in particular, involve vast complexity and uncertainty (Tambe et al., 2019). Second, human bias is assumed to be a product of this complexity, with human cognitive processes using biased shortcuts to handle complexity (Jarrahi, 2018). Third, technological change is assumed to be a naturally occurring process, impacting society and organisations (Black and Van Esch, 2020; Dwivedi et al., 2019).

Building on these assumptions, HRMOS writing often views AI through an optimist lens, offering a potential solution to organisational complexity, human bias, and changes in hiring due to increasing digitisation and online job postings. Using the term ‘digital recruiting’, Black and Van Esch (2020) trace the shift from *analog hiring* (mostly manual) to *digital recruiting* (digitised and automated practices transformed and dominated by AI), portraying AI solutions as the ‘natural’ technological response to the massive outreach that digitised hiring. Following this logic, AI is seen as the inevitable robust solution to the proliferation of job applications. Thus, we see a large consensus in HRMOS on the future being identified by ‘human-AI’ collaboration (Einola and Khoreva, 2023; Jarrahi, 2018; Kelan, 2023; Shrestha et al., 2019). Reflecting these assumptions about the future, a special focus within HRMOS literature is on conceptualising how this collaboration will unfold, and how specifically it will help to deal with uncertainty, complexity, and human bias.

For example, Jarrahi (2018: 7) develops a model of ‘human-AI’ symbiosis that can aid the ‘uncertainty’, ‘complexity’ and ‘equivocality’ of the decision-making process. Humans are seen as better equipped to make choices and decide ‘where to seek data’, while AI can ‘collect, curate, and analyse’ the chosen data. Shrestha et al. (2019) highlights the need to consider decision-making conditions, such as the ‘specificity of the decision search space’, ‘interpretability of decision-making process and outcome’, ‘decision-making speed’, and ‘replicability of outcomes’ (67–68). This approach underscores the strengths of AI, such as the speed and capacity to analyse large data, while human decision-making offers better interpretability and a more loosely defined decision space.

A strong underlying assumption in the literature is that tasks associated with recruitment and selection are particularly complex and, therefore, particularly challenging to automate

(Tambe et al., 2019). Furthermore, various tasks involved in these processes are assumed to have different levels of complexity and, therefore, require different types of expertise and knowledge. While some tasks in hiring are routine and can be easily automated (Upadhyay and Khandelwal, 2018), others require intuition and tacit knowledge of humans (Jarrahi, 2018). Bankins et al. (2024) uses the terms ‘data and AI sensitivities’ and ‘task sensitivities’ to describe the compatibility between different task characteristics with humans or AI in the particular organisational context (843).

Critical voices are also evident in the HRMOS scholarship, questioning some of these optimistic assumptions. Though more of a minority view, these perspectives rest on different assumptions: that AI raises moral and ethical considerations (Budhwar et al., 2022: 1084) and must be used as a socially responsible tool (Chang and Ke, 2024); and that AI adoption may be more complex than anticipated, given that job seekers prefer human interaction (Acikgoz et al., 2020; Newman et al., 2020) and that algorithmic decisions are perceived as fairer if humans have the final say (Newman et al., 2020; Tambe et al., 2019).

Social sciences (SS): There are many strands of research, approaches, and assumptions in this stream—rooted in sociology, psychology, and science and technology studies (STS), among others. Studies also address varied levels of analysis (e.g. micro, meso, macro). That said, one core, seemingly shared, assumption is that algorithms pose concerns for existing social inequalities, with the potential for mimicking, amplifying, or contributing to its new forms (Chen, 2023b; Lin et al., 2020).

These assumptions are linked to the social implications of technological change. While some research suggests that technological change unfolds in a way that is somewhat predictable

(i.e. we assume that computers continue to operate in a certain way), others suggest a fundamental break in the way technology is developing (Crawford et al., 2019; Elliott, 2019). Some scholars view technology as a deeply embedded socio-technical system, in which humans and technologies shape one another (Wajcman, 2017), whereas others view technology as an external force imposed on workers (Ajunwa, 2020b; Zuboff, 2019).

Building from this, a significant part of the SS scholarship problematises the technosolution framework stemming from the computational perspective to algorithmic inequality (Drage and Mackereth, 2022; Osoba et al., 2019; Rovatsos et al., 2019). Social perspectives underscore that the complexity of inequality cannot be reduced to data and its calibration, such as challenging the removal of demographic attributes from the algorithms training as a key method of ‘debiasing’ (Drage and Mackereth, 2022). Removing protected characteristics from AI training and design is seen to misfocus on the category itself rather than the systems of power that are responsible for the differential treatment of these groups.

Furthermore, understandings of equality, equity, and fairness are deeply contextual and may vary across different settings (Osoba et al., 2019; Rovatsos et al., 2019), posing a challenge for the technical and computational universalist discussion of ‘debiasing’. For instance, Osoba et al. (2019) points out that the definitions of what *procedural equity* (related to the fairness of the decision-making process) and *outcome equity* (related to the fairness of the outcome) mean different things in different contexts. Overall, the SS perspectives point to the systemic nature of inequalities, including algorithmic, that require structural, rather than individual-level and/or technical perspectives (Joyce et al., 2021).

Legal scholarship (LS): Legal scholars hold diverse views on the challenges of algorithmic hiring and decision-making. While some scholars view the foundational issues as legal (Ajunwa, 2020b; Bornstein, 2018), others emphasise technical dimensions (Blass, 2019). Many scholars adopt either implicitly or explicitly some form of socio-technical-legal approach, recognising the complexities and need for multifaceted solutions that transcend the boundaries of law. Overall, three sets of assumptions, or debates over these issues, stand out.

First, we see varied assumptions around the efficacy of existing legal frameworks to address potential problems of algorithmic bias and discrimination. Some argue they are up to the task (Nachbar, 2021), while others advocate for change (Ajunwa, 2020a). Secondly, while legal scholars are doubtful of the argument that algorithms minimise bias (Bent, 2020), many advocate for ‘technical fixes’ in data and design as the best approach for mitigating bias. Finally, some scholars assume that the nature of algorithmic bias requires more proactive, multi-party approaches (e.g. third party auditing), noting that complaint-based employment laws are insufficient (Ajunwa, 2020a; Hacker, 2018; Kim, 2020).

Building on these assumptions, legal scholars propose a variety of approaches for promoting fairness and mitigating discrimination. In the US, Ajunwa (2021a) advocates for a proactive approach that prevents bias at its source by mandating alterations to data inputs and implementing third-party certification processes involving legal, software engineering, and data science expertise. Bent (2020) explores the legalities of incorporating protected traits into algorithm design, advocating for ‘algorithmic affirmative action’ (809). Blass (2019) focuses on data protection, proposing a third-party recordkeeper system to safeguard representative data. Bornstein (2018) advocates for an approach with proactive and reactive dimensions, that involves regulating and documenting algorithmic choices before use, but also strengthening

existing law to address algorithmic disparate impacts post-occurrence. This dual approach aims to improve algorithms and legal frameworks simultaneously.

More integrated legal frameworks, particularly within the EU, offer a distinct avenue for combating algorithmic bias (Parviainen 2022). Hensler (2019) highlights the merits of the GDPR, while advocating for proactive regulation to audit algorithms for discriminatory effects. Kaminski (2018) suggests a binary approach, combining individual due process rights with collaborative governance for systemic regulation. Mann and Matzner (2019) emphasise the importance of EU regulations in protecting individual rights, especially data privacy and the right not to be subject to automated decisions. They also offer unique insights into less discussed issues of intersectional discrimination (e.g. gender $\times$ race), emergent discrimination (e.g. browsing histories), and the need for decolonial perspectives to understand how algorithmic profiling may perpetuate colonial legacies. These latter points echo Acker's (2006) emphasis on the need to centre globalisation and inequalities.

Discussion and Critical Reflections

Problematizing and identifying new questions and concerns

Building on these findings, and deepening our problematising approach, this section brings together multidisciplinary scholarship to identify novel questions and concerns about the challenges posed by AI-mediated hiring. First, in tracing multidisciplinary conversations about AI, hiring, and organisational inequalities, we note emerging asymmetries in how knowledge is being constructed and the dominance of concepts and approaches, especially from CS. Second, in examining hiring processes as foundational gatekeeping mechanisms that structure access to jobs, we note how AI is becoming intertwined with the 'inequality regimes' in organisations,

concealing inequality further within AI's opaque apparatus. Finally, we note a paradox between stated concerns over AI, regulatory lags, and contestation. Analytically, we situate these trends within critical perspectives on power asymmetries and the role of AI in both exacerbating and concealing them further. We conclude by discussing an agenda for future research and each stakeholder's imagined role in the chain of institutional and collective responsibility (Miller, 2017: 344).

Figure 2 provides a synthesis of our review and visualises the complex interplay across disciplinary domains. In the centre, the blue lines indicate interdisciplinary conversations and the interrelations of key questions and concerns across the disciplines. The dotted orange lines illustrate the chain of knowledge production across disciplines, such as asymmetries in emerging knowledge. The green lines indicate emerging strains within the chain of responsibility across the domains, such as contestation over the responsibility for the consequences of AI usage. We discuss each of those in turn in the subsections below.

[Figure 2 here]

Tracing asymmetries in emerging multidisciplinary knowledge: While all disciplinary perspectives are concerned with the potential of AI to create / reproduce inequalities in hiring processes, each brings unique assumptions and questions. CS focuses on technical solutions and adjusting data and algorithms to improve performance. HRMOS advocates for implementing AI, and advancing human-AI collaboration, to gain efficiencies and overcome perceived problems of human bias. SS scholarship brings diverse and often critical perspectives to bear, discussing how new technologies are developing within the broader social context and their potential risks for (re)producing inequalities. LS focuses on the efficacy of varied legal and regulatory solutions.

Bringing different disciplinary perspectives together allows us to identify shared concerns and, more importantly, clear asymmetries in emerging knowledge, with intellectual spillover in concepts, questions, and concerns. CS is clearly dominating discussion through expert knowledge of the subject matter: the algorithm. This is evident in the escalating presence of CS publications and perspectives within the multidisciplinary scholarship we traced, but equally important in other disciplines through the adoption of computing terminology and mathematical definitions, such as ‘algorithmic bias’, ‘algorithmic fairness’, and ‘debiasing’. A problematising approach illuminates how knowledge is embodied in language and how terminology is shaping how ‘algorithmic inequality’ is conceptualised. This tendency has several implications for knowledge generation.

First, the widespread adoption of the terms ‘bias’ and ‘debiasing’ across disciplines contributes to discussing the issue from the technosolution perspective that assumes that complex inequality can be resolved technically (with proper data, design and interpretation of the outcome) to achieve *algorithmic fairness*. Yet, designing a ‘fair’ algorithm is difficult, if not impossible, given the contradictory nature of fairness criteria (Madaio et al., 2022; Osoba et al., 2019; Rovatsos et al., 2019). Scholars note tensions between *anti-classification* (hiding the protected characteristics), *classification error parity* (ensuring that chances of positive outcomes and errors are the same), and *calibration* (ensuring that risk scores are resulting in the same the percentage outcome) (Rovatsos et al., 2019). Moreover, the meaning of fairness is deeply contextual and should be defined / operationalised in a domain-specific manner (Osoba et al., 2019). Yet, in most discussions, the CS perspective goes unchallenged.

Second, we identify a recurring tendency in CS to focus on ‘debiasing’ while looking to other disciplines for guidance on issues of justice, fairness, and equality in hiring (Chen, 2023a;

De Cremer and De Schutter, 2021), or alternatively to delegate such issues away (Belenguer, 2022). Though some computational studies in our review do engage with other disciplines, it is more often a utilitarian than conceptual exchange. Many studies seek practical suggestions on how to improve the data, the model, and / or the outcome to achieve demographic parity between different groups. Thus, we see the adoption of the four-fifths rule—a core principle in US employment law—as a widely-adopted ground truth in algorithm development and testing (Raghavan and Kim, 2023). Such ‘outsourcing’ reaffirms the centrality of a ‘technosolution’ paradigm, which in turn results in a decontextualised approach to AI-induced inequality—another observation of our review.

Decontextualisation occurs in several ways. First, the populations from which AI training data are drawn are rarely discussed, including whether these populations can be representative of other populations for which the algorithms are to be deployed. Thus, the *transferability* of data across contexts, groups, and populations—a critical issue in a globalised economy—is unclear (Yu, 2020). Second, the organisational settings for which algorithms are developed are rarely discussed. Yet, these settings vary widely, depending on industry, labour force composition, company size, location, and relevant local, national, global laws. Such contextual knowledge is actually critical for reconciling contradictory rules to achieve ‘algorithmic fairness’ (Osoba et al., 2019) and determining whether algorithmic inequalities will be minor or significant (Rovatsos et al., 2019).

Tracing how algorithms are intertwined with ‘inequality regimes’: Expectations of technical solutions to resolve issues that are broad, pervasive, and concealed, illustrate the current influence of a technosolution paradigm. Yet, within our multidisciplinary review, other

perspectives critique the expectation that algorithms can resolve issues that (re)produce inequalities via proper data and calibration as being overly simplistic and optimistic (Drage and Mackereth, 2022; Kelan, 2023; Köchling and Wehner, 2020). The framing of algorithmic tools as solutions to human bias (Lin et al., 2020; Upadhyay and Khandelwal, 2018) implies that inequalities are a problem occurring at the *individual level* (e.g. human / implicit bias), overlooking their *structural* pervasiveness and embedding in organisational contexts.

Taking a problematising approach to these issues, we argue that both human and algorithmic decision-making processes are situated within structural *inequality regimes*, reproducing relations of power in ways that are visible and invisible (Acker, 2006). Within organisations, status group inequalities are created and/or reinforced, through segregation or integration, or the granting of (un)equal access to resources or other forms of power (Ozturk and Berber, 2022; Ray, 2019; Tomaskovic-Devey et al., 2009). These inter-group dynamics reflect power relations rooted in specific economic, social, and legal contexts, at local, regional, national, and global levels. Furthermore, organisational processes are not isolated from broader power dynamics. As Pasquale (2015) argues, the proliferation of opaque, algorithmic systems reinforces these dynamics by concealing decision-making processes further, thus limiting visibility and accountability.

Specifically, the complex and opaque nature of algorithmic solutions (Elliott, 2019; Kellogg et al., 2020; Pasquale, 2015), lack of transparency (Langenkamp et al., 2020), and ability to easily incorporate explicit and implicit information about job applicants (Cruz, 2024) contributes to further invisibility of differential treatment (Acker, 2006), hidden within the apparatus of automated or augmented decision-making (Raisch and Krakowski, 2021). Building on Pasquale (2015), Ajunwa (2020a: 1724) warns that automated hiring systems ‘may become

the worst type of broker, a “*tertius bifrons*” between the employer and employee, cementing algorithmic authority over hiring processes and sustaining a new, algorithmic, ‘black box’. This renders decision-making more opaque, inequalities more concealed, and job seekers more powerless—signalling the emergence of an *algorithmically-mediated* inequality regime. This regime is characterised by extreme levels of *algorithmic invisibility*, as well as pervasive and unchallenged *algorithmic legitimacy*.

Data play a key role within algorithmically-mediated inequality regimes. Data are used to feed and train algorithms, which in turn, collect and analyse data, often without proper consent (Dattner et al., 2019). Operating as a part of an algorithmic system, data extend beyond the required application materials) to include information collected indirectly (e.g., datafied social media profiles) that may signal protected characteristics and behavioural traces. Both direct and indirect data can serve as filters for inclusion or exclusion (Cruz, 2024; Todoli-Signes, 2019). Zuboff (2019) refers to behavioural data as the new ‘gold dust’ of surveillance capitalism, offering huge predictive and economic value (Todoli-Signes, 2019). The use of online data is an increasingly contested area, where practices are difficult to regulate, not entirely illegal but potentially unethical, creating an invisible space for inequalities to occur and reproduce (Rosenblat et al., 2014; Zuboff, 2019).

The expanded use of datafied labour profiles (e.g., resume databases, LinkedIn pages, algorithmic matching scores) illustrates this concern. In traditional hiring, applicants choose how to present and narrate themselves. In algorithmic hiring, the uncontrolled access to digital data shifts control away from applicants, toward opaque algorithmic systems that collect and use digital traces in ways that are increasingly (*algorithmically*) invisible (Ajunwa, 2020a; Pasquale, 2015). As employers derive power from access to more data, candidates have less power and less

agency. Yet, attempts to legitimate current practices are made through claims of the unprecedented volume of digital job applications that, now, necessitate AI solutions, illustrating the emergence of (*algorithmic*) legitimacy that references the necessity of algorithmic solutions despite contestation and concerns. These claims of ‘algorithmic necessity’ shift the blame onto job seekers as the ones who generate massive flows of applications, while in fact, the gatekeeping power is concentrated with the employers, who can potentially benefit via an enlarged pool of applicants and data. Candidates respond to digital hiring using the platforms, but for them this means increased competition, minimised chances of being hired, and less control over the process.

Another illustration of *algorithmically-mediated inequality regimes* can be seen in the use of facial recognition to evaluate candidates. Framed as tools of efficiency, facial recognition systems measure, code, and interpret candidates’ facial expressions through opaque algorithms that are nearly impossible to question, understand, or appeal (Pasquale 2015; Ajunwa 2020a). This makes the process of evaluation *algorithmically invisible*, and generates skewed, one-sided visibility, while employers see everything and candidates see little. Codified facial expressions are what Zuboff (2019) calls ‘behavioural surplus’, predictive data for employers to define which facial cues signify the ‘ideal applicant’. This strips candidates of agency over how to present themselves in the hiring process.

Given how pervasive AI in hiring processes is becoming, it can and will intertwine with many other decision-making processes, embedding itself throughout organisational encounters and interactions. This has several implications. First, highly specialised expert knowledge is required to understand algorithmic tools, their design, and calibration (Pasquale, 2015). This knowledge is not widely shared by all stakeholders involved in the hiring process, confirming the

emergence of a new technical ‘coding elite’—stakeholders possessing exclusive knowledge on the design of algorithmic solutions (Burrell and Fourcade, 2021: 215). Second, and related to our first point, the lack of transparency and the need for elite knowledge sets pose challenges for regulation, accountability, and responsibility (Martin, 2019; Miller, 2017), a point we return to shortly.

Tracing questions of regulation, legitimacy, and accountability: A final insight from our review highlights questions about the regulation of algorithmic hiring as well as contestation over how problems of unfair treatment and accountability will be addressed. Here the LS and SS scholarship offers more critical perspectives, underscoring the need to accelerate regulation, and questioning (with some exceptions) the faith placed in purely technosolutions. That said, some legal scholars do argue that technical fixes, voluntary audit approaches, or a combination, are sufficient, with problems handled via existing legal frameworks. For instance, in the US, Sonderling et al. (2022) argue against stringent regulations, proposing that federal agencies such as the Equal Employment Opportunity Commission (EEOC) provide expert opinion to clarify the law and engage companies in voluntary compliance.

While technological advancements continue to outpace regulation, frameworks are emerging under which the multiple stakeholders involved in algorithmic hiring—employers, vendors, contractors, computer scientists, workers, unions, advocacy and civil society groups—will operate. LS offers insights on the emerging regulatory landscapes as discussed by Bradford (2023) and others (see Table 2 below). Much attention focuses on the US model, which emphasises efficiency and minimal government intervention. A smaller subset of EU-focused studies consider more ‘muscular’ models (Kim and Bodie, 2021), specifically the EU’s GDPR,

and the 2024 EU AI Act which places strong restrictions on high-stakes decision making like hiring, requiring risk impacts, mandatory pre-deployment, transparency measures, and human oversight (Kaminski, 2018; Kaminski and Malgieri, 2025). Yet, there are many questions about the impact of regulation and the potential for workarounds, such as human-led rubber stamping of algorithmic decisions to avoid illegality (Parviainen, 2022). Moreover, at the time of writing, lawmakers in the US are considering bans on AI regulation at the state level (Lima-Strong, 2025). It is also notable that other regulatory models, such as the state-centred Chinese model that Bradford (2023) outlines, and others, are not discussed, confirming the dominance of Global North perspectives. Future research and comparative analyses across systems will be important for shedding light on how distinct modes of AI may contest one another in a global AI race. Indeed, this contestation may be crucial to understanding the role of AI in reproducing inequalities on a global scale.

[Table 2 here]

Bringing the disciplines together helps pinpoint key challenges in regulating algorithmic hiring, given the myriad stages where discrimination can occur, the ways in which unfair treatment is concealed in ever more complex algorithmic ecosystems, and the limits of existing legal approaches (Osoba et al., 2019). Our review highlights diverse viewpoints, however, showing that the path ahead is not predetermined. This is illustrated in Table 2 below. In the US, Title VII of the Civil Rights Act has been a key tool for combatting employment discrimination, prohibiting *disparate treatment* and *disparate impact*. But a serious problem with the US approach is the burden of proof placed on individual complaints, prompting calls for alternatives such as proactive audits, third-party certifications, and hybrid legal approaches that blend employment and consumer protection / privacy law. Developments such as New York City's

2023 law require proactive audits to ensure fairness; employers must notify applicants when using these tools, provide other options, conduct annual audits, and publicly disclose results. Yet, loopholes—such as the inadequate independence for auditors—limit its effectiveness (Fuchs, 2023). US law also does little to aid transparency and accountability, given trade secret protections (Kim and Bodie, 2021).

Acker's (2006) ideas are useful for thinking through the current challenges of regulation and algorithmic hiring. Historically, she observes, inequalities have been successfully challenged when their visibility is high and their legitimacy is low. Today, AI threatens to further obscure this visibility, making it harder, or impossible, to identify what is (re)producing inequalities. This *algorithmic* invisibility challenges existing regulatory frameworks that rely on the paradigm of redressing the differential treatment and outcomes caused and mediated by humans. Equally important, current regulatory lag and contestation, certainly in the market-based US system, may be allowing legitimacy to grow around algorithmic hiring, positioning it as a 'business necessity', despite growing concerns.

A critical question, then, is how these *algorithmically-mediated* inequality regimes can be tempered and restrained, and what new legal paradigms are suited for this new organisational reality. Acker (2006) argues that successfully challenging inequalities in the past has required the combined force of the law, civil advocacy, and social movements outside of organisations, along with change agents inside organisations. This aligns with Zuboff's (2019) assertion that democratic oversight over surveillance capitalism necessitates mobilised collective resistance where social movements can be seen as critical agents in challenging the asymmetries embedded in and reproduced by algorithmic systems. Our review underscores the urgent need for deeper, multidisciplinary engagement by scholars in these areas, alongside allied efforts by diverse

stakeholders—e.g. workers, unions, advocacy groups, regulators—to ensure fair and accountable AI systems and their deployment.

Contributions, limitations, and future research directions

Our review makes contributions to the emerging understanding of AI, hiring processes, and organisational inequality, with implications for: 1) theoretical understandings; 2) knowledge production; and 3) practical applications. Below we outline these contributions and propose future research directions, emphasising the need for advancing theoretical understanding of AI as a part of inequality regimes, fostering interdisciplinary collaboration, and bridging gaps between academia, policy, and practice.

It is important to note that our review is limited to English-language publications, reflecting a Global North bias and under-representing perspectives from the Global South. Future research must prioritise comparative work across diverse jurisdictions and global perspectives to fully understand the complexities of algorithmic hiring. Moreover, while our hybrid methodology offers strengths in mapping and analysing emerging multidisciplinary scholarship, publication practices and quality vary across fields (e.g. CS often publishes in peer-reviewed conference proceedings rather than indexed journals). Future studies may therefore wish to employ systematic review methodologies that include explicit quality measures. Finally, while our review focuses on hiring processes as a foundational gate-keeping mechanism, other HR decision-making processes altered by AI (e.g. promotions, compensation) merit attention in future research (Manroop et al., 2024).

Theoretical understanding: By applying and extending Acker's framework, we re-conceptualise AI-transformed organisations as sites of *algorithmically-mediated* inequality regimes. These regimes are reinforced through the growing legitimacy and pervasiveness of AI-driven solutions, and their capacity to conceal inequality further within the algorithmic mechanisms that require expert knowledge to be fully understood. Framing AI as a part of 'inequality regimes' enables a deeper understanding of how AI both participates in and transforms the structure of inequality within organisations.

Future research should continue to theorise AI's role in perpetuating inequalities, advancing knowledge of AI-mediated mechanisms identified by the growing *legitimacy* and *invisibility* of high-stakes decision-making processes. Critical perspectives from varied disciplines (e.g. sociology, political economy, law, etc.) can offer valuable insights into whether AI emerges as a distinct structural mechanism that reproduces / conceals inequality, or interacts with and embeds in existing mechanisms. A particular direction for theorising should consider the multilayered relations of power beyond the organisations (Acker, 2006), such as political, economic, and international structures. Building on and extending the work of Zuboff (2019) and Pasquale (2015) could help to trace the ways in which a new economic order is spilling over and defining how AI is operating with high-stakes social institutions and its regulation. Specifically, it will be important to continue conceptualising these tendencies within organisations—particularly in AI-mediated hiring—by examining whether and how AI systems are gaining legitimacy, shifting from optional tools to imposed necessities.

Studies grounded in qualitative traditions (e.g. ethnography, interviews, historical and comparative designs) will be especially helpful for exploring the discourses, practices, and attitudes around AI, and understanding how AI may reproduce inequalities, experiences of

algorithmic inequality, and efforts to mitigate it. Specifically, ethnographic studies can trace practices related to AI through different ‘standpoints’ (Smith, 2005), such as practices of creating AI, focusing on computing science industry, practices of using AI, focusing on domains of application (e.g. workplaces, healthcare, education, etc.), and practices of regulating AI, focusing on how legal and regulatory developments. Studying these topics can elucidate political, economic, and other forms of power asymmetries, thus contributing further to conceptualising the place and role of AI in shaping inequalities.

Knowledge production: Our review highlights the key role of multidisciplinary in generating and advancing knowledge about AI applications for social inequality. Here, we can extend the concept of ‘chain of responsibility’ (Miller, 2017) that frames a collective responsibility of different stakeholders in AI application to knowledge production, highlighting that a ‘responsible’ understanding of AI requires *a chain of knowledge and expertise*. While various disciplines share a common concern for addressing social inequality, they work with distinct questions, concerns, and assumptions, as our review has shown. These disciplinary approaches typically confine analyses to domain-specific expertise and objectives, overlooking insights from other fields. For example, CS possesses the expertise on the subject matter, the algorithm, but has limited knowledge on the domains of its application in hiring, including the organisational contexts and constraints, relevant policies, and regulations. SS, HRMOS, and LS each have expert knowledge of their domains but typically lack expert knowledge to evaluate algorithmic solutions employed in their domains. These disruptions to the chain of knowledge obscure an understanding of how algorithmic solutions operate in and interact with high-stakes processes and decisions such as hiring. To address these gaps, there is a need for more meaningful

engagement between disciplinary areas with implications both for *how* to collaborate and *what* to collaborate on.

In terms of *how* to collaborate, fostering more meaningful engagement between disciplinary areas is essential for achieving a true post-disciplinary approach to AI, one that is human grounded, socially embedded, and that challenges the technosolution paradigm circulating around AI use and mitigation of its impact. Currently, CS dominates much of the current discourse. However, insights from SS, LS, and, HRMOS studies are critical to understanding the organisational, institutional, and structural forces in place. Collaborative efforts across these fields can lead to a more comprehensive understanding of what ethical and responsible AI can look like, fostering solutions that incorporate both technical and contextual, domain-specific, dimensions. Specific mechanisms that foster interdisciplinary research, such as joint conferences, funding initiatives, research networks, or special issue calls can facilitate collaboration across disciplines and sectors, allowing vital exchange between academics, algorithm designers, policy creators and legal experts.

In terms of *what* to collaborate on, our review highlights many pressing issues. First, the foundational concepts of algorithmic (in)equality, fairness, ethics, and transparency require rethinking, as current definitions often fail to address entrenched systemic inequalities and contextual characteristics. Second, a multi-level approach—spanning micro, meso, and macro—is crucial for developing frameworks that prioritise structural approaches to inequality in the organisational application of AI, with more awareness of the specific level at which research is conducted. Third, research on AI and hiring demands stronger contextualisation to account for differences across industries, company size, and workforce compositions. For example, research could focus on implications of AI solutions for historic, complex, and intersectional inequalities.

Fourth, governance of algorithmic hiring systems requires multidisciplinary research to ensure equity, feasibility, and accountability. Initiatives such as the EU AI Act exemplify the multidisciplinary and cross-sectoral approach to developing AI governance frameworks (Rotenberg and Kyriakides, 2025). We suggest a stronger scholarly focus on tracing how the *chain of knowledge* operates within such initiatives, examining how ideas, assumptions, and approaches cooperate.

Practical application: Building on insights into the *chain of knowledge*, another priority is to operationalise *the chain of responsibility* at a practical level. As our review shows, algorithmic hiring systems can obscure structural inequalities while deflecting accountability across multiple stakeholders. Because challenges cannot be solved by a single stakeholder, we advocate for a collaborative agenda involving employers, organisations, algorithm designers, unions, regulators, and policy makers, with each assuming and maintaining their unique responsibility in the chain. In terms of concrete steps, employers can play a critical role through participatory design practices and voluntary audits to ensure AI systems are transparent and fair. They can also build capacity and awareness on responsible AI use by fostering partnerships between the HR professionals, algorithm designers, and regulators, to evaluate real-world AI applications, identify barriers to fairness, and embed these insights into their own hiring practices. Likewise, legal practitioners and scholars can work with policy makers to craft regulations that respond to the evolving challenges that AI presents. Interdisciplinary resources can bridge gaps between legal, policy, and organisational domains. Collaborative efforts must be tailored to specific contexts (e.g. industry, country), recognising that equity and fairness are deeply contextual. One example is *The AI Policy Sourcebook* (Rotenberg and Kyriakides, 2025), a comprehensive

handbook of AI policy frameworks from around the world, that provides a valuable practical tool for coordinated action.

It is important to recognise that power amongst diverse stakeholders is uneven and not all actors will choose, or be able, to act. Some employers may adopt responsible AI practices, others may not; regulators may intervene, or remain inactive. Civil society groups (such as Algorithmic Justice League, AI Now Institute, Upturn) will be critical in ensuring responsible AI practices, by engaging in public advocacy, submitting regulatory briefs, and educating job seekers about AI in hiring, while holding employers and developers accountable. As Zuboff (2019) and Pasquale (2015), and earlier Acker (2006), argue, their work is vital for making structures of inequality visible and creating pressure for change.

Conclusions

Prompted by the growing use of AI in hiring, and concerns over organisational inequalities, this review combines scoping and problematising approaches to assess emerging knowledge. We examine four broad disciplines—computing science, human resource management, social sciences, and law—that are largely siloed, tracing their key concerns and assumptions. Our review shows that 1) computing science’s dominance is driving attention to technosolutions, with the language of ‘bias’ spilling into other fields; 2) framing AI as a solution to human bias diverts attention from structural inequalities that are evolving into the *algorithmically-mediated* inequality regimes; and 3) regulatory responses are lagging, uneven, and contested, with solutions shaped by Global North perspectives. Our future research agenda emphasises the need for deeper theoretical insights; an interdisciplinary *chain of knowledge* creation; and operationalising the *chain of responsibility* for practical solutions. Building on Acker (2006), we

stress that inequalities are most effectively addressed when visible and illegitimate—the opposite of what we currently see in the growing use of AI in organisational hiring.

References

- Acikgoz Y, Davison KH, Compagnone M and Laske M (2020) Justice perceptions of artificial intelligence in selection. *International Journal of Selection and Assessment* 28(4): 399-416. [*]
- Acker J (2006) Inequality regimes: Gender, class and race in organizations. *Gender & Society* 20(4): 441-464.
- Aizenberg E and Van Den Hoven J (2020) Designing for human rights in AI. *Big Data & Society*, 7(2): 1-19. [*]
- Ajunwa I (2020a) The paradox of automation as anti-bias intervention. *Cardozo Law Review* 41: 1671 - 1742. [*]
- Ajunwa I (2020b) The “black box” at work. *Big Data & Society* 7(2): 1-6. [*]
- Ajunwa I (2021a) An auditing imperative for automated hiring systems. *Harvard Journal of Law & Technology* 34(2): 621-699. [*]
- Ajunwa I (2021b) Automated video interviewing as the new phrenology. *Berkeley Technology Law Journal* 36: 1173-1226. [*]
- Ajunwa I and Greene D (2019) Platforms at work: Automated hiring platforms and other new intermediaries in the organization of work. In: Vallas SP, Kovalainen A (eds) *Work and Labor in the Digital Age: Research in the Sociology of Work, Vol 33*. Leeds: Emerald Publishing Limited: 61-91. [*]

- Allal-Chérif O, Yela Aránega A and Castaño Sánchez R (2021) Intelligent recruitment: How to identify, select, and retain talents from around the world using artificial intelligence. *Technological Forecasting and Social Change* 169: 1-11. [*]
- Alvesson M and Sandberg J (2020) The problematizing review: A counterpoint to Elsbach and Van Knippenberg's argument for integrative reviews. *Journal of Management Studies* 57(6): 1290-1304.
- Alvesson M and Skoldberg K (2018) *Reflexive Methodology: New Vistas for Qualitative Research*. London: Sage Publications.
- Amis JM, Mair J and Munir KA (2020) The organizational reproduction of inequality. *Academy of Management Annals* 14(1): 195-230.
- Amershi S, Cakmak M, Knox WB and Kulesza T (2014) Power to the people: The role of humans in interactive machine learning. *AI Magazine* 35(4): 105-120. [*]
- Arksey H and O'Malley L (2005) Scoping studies: Towards a methodological framework. *International Journal of Social Research Methodology* 8: 19–32.
- Baer T (2019) *Understand, Manage, and Prevent Algorithmic Bias: A Guide for Business Users and Data Scientists*. New York, NY: Apress. [*]
- Bankins S, Ocampo AC, Marrone M, Restubog SLD and Woo SE (2024) A multilevel review of artificial intelligence in organizations: Implications for organizational behavior research and practice. *Journal of Organizational Behavior* 45(2): 159-182. [*]
- Basu S, Majumdar B, Mukherjee K, Munjal S and Palaksha C (2023) Artificial intelligence-HRM interactions and outcomes: A systematic review and casual configurational explanation. *Human Resource Management Review* 33(1): 1-16. [*]

- Belenguer L (2022) AI bias: exploring discriminatory algorithmic decision-making models and the application of possible machine-centric solutions adapted from the pharmaceutical industry. *AI and Ethics* 2(4): 771-787. [*]
- Bent JR (2020) Is algorithmic affirmative action legal? *Georgetown Law Journal* 108(4): 803-854. [*]
- Berkelaar BL (2017) Different ways new information technologies influence conventional organizational practices and employment relationships: The case of cybervetting for personnel selection. *Human Relations* 70(9): 1115-1140.
- Bills DB, Di Stasio V and Gërkhani, K (2017). The demand side of hiring: Employers in the labor market. *Annual Review of Sociology* 43(1): 291-310.
- Birhane A and Cummins F (2019) Algorithmic injustices: Towards a relational ethics. *arXiv preprint arXiv:1912.07376*: 1-4. [*]
- Black SJ and Van Esch P (2020) AI-enabled recruiting: What it is and how a manager should use it? *Business Horizons* 63(2): 215-226. [*]
- Blass J (2019) Algorithmic advertising discrimination. *Northwestern University Law Review*, 114(2): 415-467. [*]
- Bornstein, S (2018) Antidiscriminatory algorithms. *Alabama Law Review* 70(2): 519-572. [*]
- Bradford, A (2023) *Digital Empires: The Global Battle to Regulate Technology*. Oxford University Press.
- Budhwar P, Malik A, De Silva MTT and Thedushika P (2022) Artificial intelligence--Challenges and opportunities for international HRM: A review and research agenda. *The International Journal of Human Resource Management* 33(6): 1065-1097. [*]

- Burrell J and Fourcade M (2021) The society of algorithms. *Annual Review of Sociology* 47(1): 213–237.
- Chang YL and Ke J (2024) Socially responsible artificial intelligence empowered people analytics: A novel framework towards sustainability. *Human Resource Development Review* 23(1): 88-120. [*]
- Chen Z (2023a) Collaboration among recruiters and artificial intelligence: removing human prejudices in employment. *Cognition, Technology & Work* 25(1): 135-149. [*]
- Chen Z (2023b) Ethics and discrimination in artificial intelligence-enabled recruitment practices. *Humanities and Social Sciences Communications* 10(1): 1-12. [*]
- Crawford K, Dobbe R, Dryer T, Fried G, Green B, Kaziunas E, Kak A, Mathur V, McElroy E, Nill Sánchez A, Raji D, Lisi Rankin J, Richardson R, Schultz J, West S, and Whittaker M (2019) *AI Now 2019 Report*. New York: AI Now Institute. [*]
- Cruz IF (2024) How process experts enable and constrain fairness in AI-driven hiring. *International Journal of Communication* 18: 656-676. [*]
- Dattner B, Chamorro-Premuzic T, Buchband R and Schettler L (2019) The legal and ethical implications of using AI in hiring. *Harvard Business Review*: 1-7. [*]
- Daudt HML, Van Mossel C and Scott SJ (2013) Enhancing the scoping study methodology: A large, inter-professional team's experience with Arksey and O'Malley's framework. *BMC Medical Research Methodology* 13(48): 1-9.
- De Alford G, Hayden SK, Wittlin N and Atwood A (2020) Reducing age bias in machine learning: An algorithmic approach. *SMU Data Science Review* 3(2): 1-19. [*]
- De Cremer D and De Schutter L (2021) How to use algorithmic decision-making to promote inclusiveness in organisations. *AI and Ethics* 1: 563-567. [*]

- DiPrete TA and Eirich GM (2006). Cumulative advantage as a mechanism for inequality: A review of theoretical and empirical developments. *Annual Review of Sociology* 32(1): 271-297.
- Drage E and Mackereth K (2022) Does AI debias recruitment? Race, gender, and AI's "eradication of difference". *Philosophy & Technology* 35(89): 1-25. [*]
- Dwivedi YK, Hughes L, Ismagilova E, Aarts G, Coombs C, Crick T, Duan Y ... and Williams MD (2019) Artificial intelligence (AI): Multidisciplinary perspectives on emerging challenges, opportunities, and agenda for research, practice and policy. *International Journal of Information Management* 57: 1-47. [*]
- Einola K and Khoreva V (2023) Best friend or broken tool? Exploring the co-existence of humans and artificial intelligence in the workplace environment. *Human Resource Management* 62(1): 117-135. [*]
- Elliott A (2019) *The Culture of AI: Everyday Life and the Digital Revolution*. London: Routledge.
- Favaretto M, De Clercq E and Elger BS (2019) Big Data and discrimination: Perils, promises and solutions: A systematic review. *Journal of Big Data* 6(1): 1-12. [*]
- Friedman GD and McCarthy T (2020) Employment law red flags in the use of artificial intelligence in hiring. *Business Law Today*: 1-14. [*]
- Fuchs L (2023) Hired by machine: Can a New York City law enforce algorithmic fairness in hiring practices? *Fordham Journal of Corporate and Financial Law* 28(1): 185-222. [*]
- Gelles R, McElfresh D and Mittu A (2018) Project Report: Perceptions of AI in Hiring. Paper presented at the University of Maryland. College Park, MD. [*]

- Geyik SC, Ambler S and Kenthapadi K (2019) Fairness-aware ranking in search & recommendation systems with application to linkedin talent search. Paper presented at the ACM Conference on Knowledge Discovery & Data Mining, Anchorage, AK. [*]
- Glymour B and Herington J (2019) Measuring the biases that matter: The ethical and casual foundations for measures of fairness in algorithms. Paper presented at the Proceedings of the ACM Conference on Fairness, Accountability, and Transparency - FAT'19, Atlanta, GA. [*]
- Gonen H and Goldberg Y (2019) Lipstick on a pig: Debiasing methods cover up systemic gender biases in word embeddings but do not remove them. *arXiv:1903.03862*: 1-6. [*]
- Hacker P (2018) Teaching fairness to artificial intelligence: Existing and novel strategies against algorithmic discrimination under EU law. *Common Market Law Review* 55(4): 1143-1185. [*]
- Hensler J (2019) Algorithms as allies: regulating new technologies in the fight for workplace equality. *Temple International & Comparative Law Journal* 34(1): 31-59. [*]
- Huang MH and Rust RT (2018) Artificial intelligence in service. *Journal of Service Research* 21(2): 155-172. [*]
- Jarrahi MH (2018) Artificial intelligence and the future of work: Human-AI symbiosis in organizational decision making. *Business Horizons* 61(4): 577–586. [*]
- Joyce K, Smith-Doerr L, Alegria S, Bell S, Cruz T, Hoffman SG, Noble SU and Shestakofsky B (2021) Toward a sociology of artificial intelligence: A call for research on inequalities and structural change. *Socius: Sociological Research for a Dynamic World* 7: 1-11. [*]
- Kaminski, ME (2018) Binary governance: Lessons from the GDPR's approach to algorithmic accountability. *Southern California Law Review* 92(6): 1529-1616. [*]

- Kaminski, ME and Malgieri G (2025) Impacted stakeholder participation in AI and data governance. *Yale Journal of Law and Technology* 27(1): 247-326.
- Kelan EK (2023) Algorithmic inclusion: Shaping the predictive algorithms of artificial intelligence in hiring. *Human Resource Management Journal*: 1-14. [*]
- Kellogg KC, Valentine MA and Christin A (2020) Algorithms at work: The new contested terrain of control. *Academy of Management Annals* 14(1): 366-410. [*]
- Kim PT (2020) Manipulating opportunity. *Virginia Law Review* 106(4): 867-936. [*]
- Kim PT and Bodie MT (2021) Artificial intelligence and the challenges of workplace discrimination and privacy. *ABA Journal of Labor and Employment Law* 35: 289-315. [*]
- Köchling A and Wehner MC (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*: 13(3): 795-848. [*]
- Kordzadeh N and Ghasemaghahi M (2022) Algorithmic bias: review, synthesis, and future research directions. *European Journal of Information Systems* 31(3): 388-409. [*]
- Kuhlman C, Jackson L and Chunara R (2020) No computation without representation: Avoiding data and algorithm bias through diversity. *arXiv:2002.11836 [cs]*: 1- 10. [*]
- Langenkamp M, Costa A, and Cheung C (2020) Hiring fairly in the age of algorithms. *arXiv preprint arXiv:2004.07132*: 1-33. [*]
- Lee M, Kusbit D, Metsky E and Dabbish L (2015) Working with machines: The impact of algorithmic and data-driven management on human workers. Paper presented at Proceedings of the 33rd Annual ACM Conference on Human Factors in Computing Systems - CHI '15, Seoul, Republic of Korea. [*]

- Lee MK (2018) Understanding perception of algorithmic decisions: Fairness, trust and emotion in response to algorithmic management. *Big Data & Society* 5(1): 1-16. [*]
- Lepri B, Oliver N, Letouzé E, Pentland A and Vinck P (2018) Fair, transparent, and accountable algorithmic decision-making processes. *Philosophy and Technology* 31(4): 611-627. [*]
- Lima-Strong, J (2025) US House passes 10-year moratorium on State AI laws. *Tech Policy Press*. Available at: <https://www.techpolicy.press/us-house-passes-10year-moratorium-on-state-ai-laws/>
- Lin YT, Hung TW and Huang LTL (2020) Engineering equity: How AI can help reduce the harm of implicit bias. *Philosophy & Technology* 34: 65-90. [*]
- Madaio M, Egede L, Subramanyam H, Wortman Vaughan J and Wallach H (2022) Assessing the fairness of AI systems: AI practitioners' processes, challenges, and needs for support. *Proceedings of the ACM on Human-Computer Interaction* 6(CSCW1): 1-26. [*]
- Mann M and Matzner T (2019) Challenging algorithmic profiling: The limits of data protection and anti-discrimination in responding to emergent discrimination. *Big Data & Society* 6(2): 1-11. [*]
- Manroop L, Malik A and Milner M (2024). The ethical implications of big data in human resource management. *Human Resource Management Review* 34(2): 101012.
- Martin K (2019) Ethical implications and accountability of algorithms. *Journal of Business Ethics* 160(4): 835-850. [*]
- Mauthner NA and Doucet A (2008) 'Knowledge once divided can be hard to put together again': An epistemological critique of collaborative and team-based research practices. *Sociology* 42: 971-985.

- Mays N, Roberts E and Popay J (2001) Synthesising research evidence. In: Fulop A, Allen P, Clarke A, Black N (eds) *Studying the Organisation and Delivery of Health Services: Research Methods*. London: Routledge, 188-220.
- Miller, S (2017) Institutional responsibility. In: Jankovic M and Ludwig K (eds) *The Routledge Handbook of Collective Intentionality*. New York, NY: Routledge, 338-348.
- Nachbar TB (2021) Algorithmic fairness, algorithmic discrimination. *Florida State University Law Review* 48: 509-558. [*]
- Nawaz N (2019) How far have we come with the study of artificial intelligence for the recruitment process? *Int. J. Sci. Technol. Res* 8: 488–493. [*]
- Newman D, Fast N and Harmon D (2020) When eliminating bias isn't fair: Algorithmic reductionism and procedural justice in human resource decisions. *Organizational Behaviour and Human Decision Processes* 160: 149-167. [*]
- Osoba O, Boudreaux B, Saunders J, Irwin J, Mueller P and Cherney S (2019) *Algorithmic Equity: A Framework for Social Applications*. Santa Monica, Rand Corporation. [*]
- Ozturk MB and Berber A (2022) Racialised professionals' experiences of selective incivility in organisations: A multi-level analysis of subtle racism. *Human Relations* 75(2): 213-239.
- Parviainen H (2022) Can algorithmic recruitment systems lawfully utilise automated decision-making in the EU? *European Labour Law Journal* 13(2): 225-248. [*]
- Pasquale F (2015) *The Black Box Society: The Secret Algorithms that Control Money and Information*. Cambridge, MA: Harvard University Press.
- Raghavan M, Barcos S, Kleinberg J and Levy K (2020) Mitigating bias in algorithmic hiring: evaluating claims and practices. *Proceedings of the ACM Conference on Fairness, Accountability, and Transparency*. [arXiv:1906.09208v3](https://arxiv.org/abs/1906.09208v3): 1-24. [*]

- Raghavan M and Kim PT (2023) Limitations of the “four-fifths rule” and statistical parity tests for measuring fairness. *Georgetown Law Technology Review* 8(1): 93-123. [*]
- Ray V (2019) A theory of racialized organizations. *American Sociological Review* 84(1): 26-53.
- Raisch S and Krakowski S (2021) Artificial intelligence and management: The automation-augmentation paradox. *Academy of Management Review* 46(1): 192-210. [*]
- Rivera LA (2020) Employer decision making. *Annual Review of Sociology*, 46(1), 215-232.
- Roemmich K, Rosenberg T, Fan S and Andalibi N (2023) Values in emotion artificial intelligence hiring services: Technosolutions to organizational problems. *Proceedings of the ACM on Human-Computer Interaction*, 7(CSCW1), 1-28. [*]
- Rosenblat A, Kneese T and Boyd D (2014) Networked employment discrimination. *Open Society Foundations' Future of Work Commissioned Research Papers 2014*:1-17.[*]
- Rotenberg M and Kyriakides E (2025) *The AI Policy Sourcebook*. Washington, DC: Centre for AI and Digital Policy.
- Rovatsos M, Mittelstadt B and Koene A (2019) Landscape summary: Bias in algorithmic decision-making: What is bias in algorithmic decision-making, how can we identify it, and how can we mitigate it? Centre for Data Ethics and Innovation. [*]
- Shrestha YR, Ben-Menahem SM and von Krogh G (2019) Organizational decision-making structures in the age of artificial intelligence. *California Management Review* 61(4): 66-83. [*]
- Russell SJ and Norvig P (2021) *Artificial Intelligence: A Modern Approach*. (4th ed.). Hoboken NJ: Pearson.
- Smith, D. E. (2005). *Institutional Ethnography: A Sociology for People*. Walnut Creek, CA Altamira Press.

- SHRM (Society for Human Resource Management) (2022) Automation and AI in HR. Available at: <https://www.shrm.org/content/dam/en/shrm/topics-tools/news/technology/SHRM-2022-Automation-AI-Research.pdf>
- Sonderling KE, Kelley BJ and Casimir L (2022) The promise and peril: Artificial intelligence and employment discrimination. *University of Miami Law Review* 77(1): 1-87. [*]
- Tambe P, Cappelli P and Yakubovich V (2019) Artificial intelligence in human resources management. *California Management Review* 61(4): 15-42. [*]
- Todoli-Signes A (2019) Algorithms, artificial intelligence and automated decisions concerning workers and the risks of discrimination. *Transfer: European Review of Labour and Research* 25(4): 564-481. [*]
- Tomaskovic-Devey D, Avent-Holt D, Zimmer C and Harding S (2009) The categorical generation of organizational inequality: A comparative test of Tilly's durable inequality. *Research in Social Stratification and Mobility* 27(3): 128-142.
- Upadhyay AK and Khandelwal K (2018) Applying artificial intelligence: Implications for recruitment. *Strategic HR Review* 17: 255–258. [*]
- Wajcman J (2017) Automation: Is it really different this time? *The British Journal of Sociology* 68: 119–127. [*]
- Weber M (1978) Status groups and classes. In: Roth G and Wittich C (eds) *Economy and Society*. Vol 1. Berkeley, CA: University of California Press: 302–306.
- Weiskopf R and Hansen HK (2023) Algorithmic governmentality and the space of ethics: Examples from ‘People Analytics’. *Human Relations* 76(3): 483-506. [*]
- Yu P (2020) The algorithmic divide and equality in the age of artificial intelligence. *Florida Law Review* 72: 331-389. [*]

Zuboff S (2019) *The Age of Surveillance Capitalism: The Fight for a Human Future at the New Frontier of Power*. New York, NY: Public Affairs. [*]

Figure 1: Map of the hybrid review approach

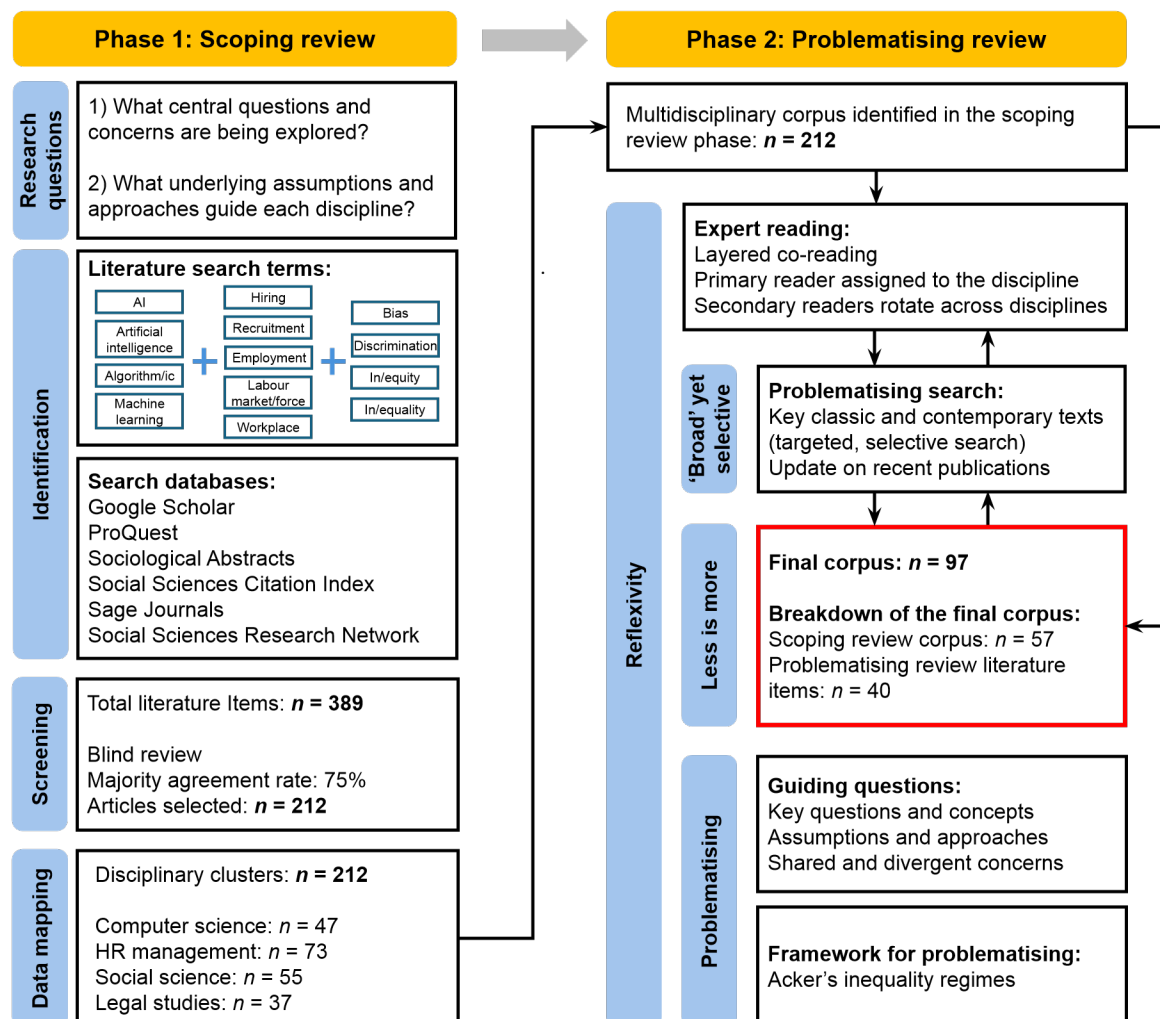


Table 1: Summary of key findings and insights

Discipline	Question 1: Key questions examined	Question 2: Underlying assumptions	Integration: Key insights & future research
CS	• Developing AI for efficiency and fairness in hiring. • Measuring and mitigating bias (e.g. demographic parity). • Debiasing datasets & algorithms. • Procedural justice in AI development and deployment.	• Bias originates from data and algorithms. • Technical solutions can address known biases but may fail to eliminate latent bias.	Key insights: • Emerging multidisciplinary discussion but with asymmetries in how key concerns are conceptualised. • Clear, heightened, potential for AI to conceal inequalities in hiring processes. • Contestation and lag over regulation and the ‘chain of responsibility’.
HRMOS	• Using AI to enhance efficiency, speed, and fairness in hiring. • AI’s role in automating vs. augmenting human decision-making. • Optimal collaboration between AI and humans in hiring processes.	• Organisational complexity and human bias necessitate AI. • AI-human collaboration is the future of hiring processes. • Technological change is a naturally occurring process.	
SS	• AI’s impact on social inequalities and discrimination. • Surveillance and data use in hiring • How AI is deployed in organisations • Regulatory frameworks for responsible AI use.	• Inequalities are systemic and deeply contextual. • Technical solutions alone cannot address structural inequalities.	
LS	• Legal frameworks for addressing AI bias and discrimination. • Adequacy of existing laws and need for proactive regulation. • Accountability and fairness in algorithmic decision-making.	• Existing legal frameworks may be inadequate to address algorithmic bias. • Proactive, multi-party approaches are required.	

Figure 2: Review synthesis - Interconnected dynamics across disciplinary domains

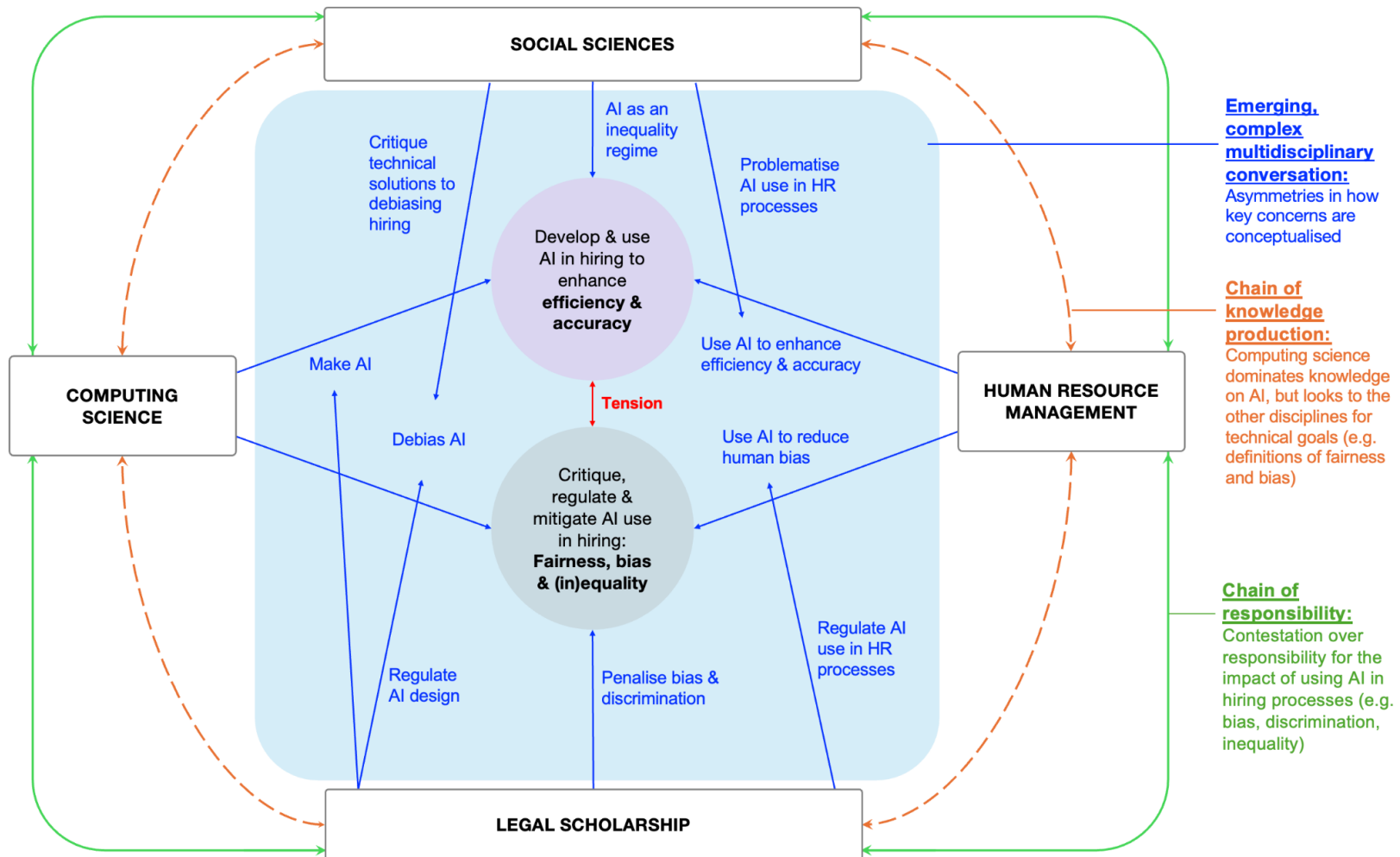


Table 2: Summary of current regulatory approaches

Dimension	United States	Europe (EU/UK)
Regulatory approach	Market-driven, complaint-based: Discrimination is addressed through litigation. Burden is placed on individuals to prove and contest discrimination.	Rights-based, preventative: Comprehensive regulation offers preventative safeguards. Responsibility for fair treatment lies with employers and regulators.
Key legislation and definitions	Anti-discrimination statutes (e.g., Title VII, EEOC) with some state / local AI-specific measures (e.g., NYC bias audits). Bias and un/fairness are established through measurable disparate treatment and impact (e.g., four-fifths rule).	The 2024 EU AI Act regulates ‘high-risk’ (e.g., hiring) decision-making. It establishes broad definitions of fairness, and mandates risk impacts, data protection, mandatory pre-deployment, transparency measures, and human oversight of decisions.
Enforcement mechanisms	Reactive, with a few exceptions (e.g., NYC audits). Investigations and remedies depend on specific complaints, and vary across jurisdictions. Without legal challenges, discrimination is left unaddressed.	Proactive. Regulators can demand evidence of compliance, impose conformity assessments, and levy significant fines. Current regulations seek to prevent harm but enforcement may vary across member states.
Consequences for applicants	Limited rights: Applicants are typically unaware when algorithmic tools are used. Onus is on applicants and advocacy groups for litigation and complaints.	Stronger procedural rights: notification of AI use, access to human review, avenues for appeal. Applicants can contest unfair practices directly.
Organizational interventions	Voluntary in most jurisdictions. Proactive organisations can embrace privacy and data protection protocols, internal audits and/or third-party auditing, model documentation, and cross-functional fairness committees to identify and mitigate risk and ensure fair treatment.	Mandatory. Regulations require organisations to carry out risk assessments, data governance, and bias testing for high-risk AI. Organisations can also institutionalise ethics by design, participatory design approaches, and continuous post-deployment monitoring. These measures translate rights to fairness, transparency, and accountability into everyday HR practices.

Supplemental material A: Scoping review – coding scheme, charting the data

Code category	Code name	Description
Bibliography codes	Author	Author names
	Title	Title of publication
	Year	Year of publication
	Journal	Name of journal, conference proceeding or publication venue
	Discipline/cluster	Academic field (e.g., social sciences, computer science)
	Type of manuscript	Type of publication (empirical, review, etc.)
	Methodology	What is the main methodology(ies) employed in the study?
Thematic codes	Key argument	What are the key arguments in this article?
	Key findings	What are the key findings/results?
	Limitations	What are the key limitations/gaps?
	Future research	What directions for future research are identified?
	Additional insights	Is there anything else important about this article?
Conceptual codes	Theoretical framework	What theory/approach is used to frame the study?
	Key concepts	What are the key concepts used/developed in the study?
	Bias terminology	Is the study using the term ‘bias’ (yes/no)?
	Bias definition	How is bias defined in the article (if applicable)?
	Types of bias	What are the types of bias discussed in the paper (e.g., human bias, implicit bias, etc.)?
	AI terminology	Is the study using the term ‘AI’, ‘algorithm’ or else (yes/no)?
	AI definition	How is AI defined in the article (if applicable)?
	Types of AI	Type of AI discussed in the paper (e.g., machine learning, automated hiring, etc.)
	Fairness/ethics terminology	Is the study using the term ‘fairness’, ‘ethics’ or else (yes/no)?
	Fairness/ethics definition	How are ‘fairness’, ‘ethics’ or else defined in the article (if applicable)?

Supplemental material B: Summary of the reviewed articles in the final corpus.

Publication	Cluster	Key foci	Key concepts ¹
Acikgoz et al. (2020)	HRMOS	Perceptions of AI	AI, ‘AI life cycle’, fairness, justice, perceptions, trust, selection decisions
Adams-Prassl (2019)	LS	Algorithmic management	Algorithmic management, control, data protection, discrimination, fairness, Polyani’s paradox, privacy, responsibility
Aizenberg and Van Den Hoven (2020)	SS	Human rights implications of AI	AI, accountability, fairness, human rights, privacy
Ajunwa (2020a)	LS	Implications of algorithmic hiring for inequality and discrimination	Algorithmic hiring, anti-bias, automation, bias, cultural fit, data, disparate treatment and disparate impact, discrimination, efficiency, fairness, law, platform authoritarianism
Ajunwa (2020b)	SS	Implications of algorithmic hiring for inequality and discrimination	Algorithmic hiring, algorithmic invisibility, black box, ‘data laundering’, discrimination, surveillance
Ajunwa (2021a)	LS	Implications of algorithmic hiring for inequality and discrimination	AI, algorithmic hiring, bias, discrimination, fairness, legal accountability, regulation, technical transparency
Ajunwa (2021b)	LS	Ethical and legal implications of automated video interviewing (AVI)	Anti-discrimination law, automated hiring, automated video interviewing (AVI), bias, discrimination, fairness, hiring, privacy, regulation
Ajunwa and Greene (2019)	SS	Automated hiring platforms	Automated hiring platforms, digital intermediaries, platform authoritarianism, management, work platforms
Allal-Chérif et al. (2021)	HRMOS	Digital technologies and recruitment	AI, automation, decision-making process, digital technologies, e-recruitment, efficiency, human bias, human oversight
Amershi et al. (2014)	CS	Roles of humans in interactive machine learning	Humans, intelligent systems, interactive machine learning, systems, users
Andrews and Bucher (2022)	LS	Implications of algorithmic hiring for inequality and discrimination	AI, algorithm training, automated resume screening, video interviewing, video games, bias, data, discrimination, gender inequality, hiring
Baer (2019)	CS	Measuring and mitigating algorithmic bias	AI, algorithmic bias detection, debiasing, mitigation, prevention
Bankins (2021)	HRMOS	Ethical use of AI in human resource management (HRM)	AI, accountability, bias, decision-making process, fairness, ethical AI, HRM, task-technology fit
Bankins et al. (2024)	HRMOS	Algorithmic management	AI, algorithmic management, ethics, employee interactions, human-AI collaboration, perceptions, platform-based work, work

Basu et al. (2023)	HRMOS	Implications of AI for HRM	AI, AI-HRM interaction, bias, efficiency, employee engagement, HRM, human oversight, governing, organisational performance
Belenguer (2022)	CS	Machine-centric solutions to algorithmic bias	AI, algorithmic bias, decision-making, fairness, human-centric solutions, machine-centric solutions, testing, transparency, validation
Benbaya et al. (2020)	HRMOS	Implications of AI in organisational decision-making	AI, data, decision-making, explainability, ethics, organisations, transparency
Bent (2020)	LS	Implications of algorithmic affirmative action	Algorithmic affirmative action, anti-discrimination law, bias, fairness, governmental actors/interest/use, discrimination
Birhane and Cummins (2019)	CS	Algorithmic bias	Algorithmic decision-making, injustice, automated systems, bias, efficiency, relational ethics/justice, technical solutions
Black and van Esch (2020)	HRMOS	Implications of AI for hiring	AI-enabled recruiting, bias, data-driven insights, decision-making, digital recruiting technology, efficiency, human resources
Blass (2019)	LS	Algorithmic advertising discrimination	Algorithmic accountability, algorithmic advertising, data, discrimination, legal implications, social media, transparency
Bloch-Wehba (2022)	LS	Algorithmic governance and regulation	Algorithmic governance, algorithmic audits, bias, discrimination, oversight, regulation
Borges et al (2021)	HRMOS	Implications of AI in organisational decision-making	AI, advantages, automation, business strategy, decision-making
Bornstein (2018)	LS	Algorithmic discrimination	AI, algorithmic discrimination, anti-discriminatory algorithms, anti-stereotyping, bias, decision-making processes, disparate treatment, regulation
Budhwar et al. (2022)	HRMOS	Implications of AI in transforming HRM practices	AI, AI-based applications, data privacy, decision-making processes, ethical, legal and moral concerns, global context, HRM, talent management
Chang and Ke (2024)	HRMOS	Implications of AI in transforming HRM practices	AI, equality, ethics, fairness, human resource management (HRM), inclusivity, organisations, people analytics (PA), socially responsible AI (SRAI)
Chen (2023a)	CS	AI solutions for mitigating human bias in recruitment	AI, employment, human bias, human judgment, human prejudice, recruitment, talent acquisition
Chen (2023b)	SS	Implications of algorithmic hiring for inequality and discrimination	AI-enabled recruitment, algorithmic bias and discrimination, algorithm designers, data, ethical governance, external oversight, perceptions
Cofone (2018)	LS	Algorithmic discrimination with a focus on data	Algorithmic discrimination, anti-discrimination law, data governance, disparate treatment, fairness, information, privacy, regulation, transparency

Crawford et al. (2019)	SS	Social and ethical implications of AI technologies	Algorithmic accountability, AI, bias, corporate ethics, discrimination, ethics, law, policy, power, transparency
Cruz (2024)	SS	Role of experts in defining concepts of fairness in algorithmic hiring	AI-based hiring decisions, experts, fairness, hiring, organisations
Dattner et al. (2019)	HRMOS	Legal and ethical implications of using AI in hiring	AI, bias, data, decision-making, ethical and legal implications, hiring, privacy
de Alford et al. (2020)	CS	Algorithmic bias in machine learning	AI, adversarial learning, age bias, demographic parity, ethics, fairness, machine learning, model accuracy, prediction
De Cremer and De Schutter (2021)	CS	Algorithmic decision-making	Algorithmic decision-making, data, diversity, human bias, inclusiveness, organisations, recruitment
Drage and Mackereth (2022)	SS	AI implications for hiring/recruitment bias	AI, bias, debiasing, discrimination, eradication of difference, gender, HR, inequality, race, recruitment
Dries et al. (2023)	HRMOS	AI and the future of work	AI, automation, future of work, job, robot, technology, transformation, workplace
Dwivedi et al. (2019)	HRMOS	Implications of AI for organisations, industry, society	AI, accountability, augmentation, autonomous intelligence systems, bias, business, ethics, governance, implications, machine-learning, privacy, productivity, replacing human tasks, regulation, responsibility, safety
Einola and Khoreva (2023)	HRMOS	Human-AI collaboration in HRM	AI, augmentation, automation, human-AI collaboration/co-existence, HRM, organisations
Elliott (2019)	SS	Implications of AI for culture and society	AI, automation, big data, culture, digital era, employment, humans, machines, social interactions, self and private life, surveillance, work
Favaretto et al. (2019)	CS	Implications of big data for inequalities	Algorithmic processing, bias, big data, data mining, discrimination, disparity, fairness, inequality
Fernández-Macías et al. (2018)	CS	AI and the future of work	AI, alteration, automation, autonomy, generality, occupations, organisations, skills, socio-economic and computational perspectives, task
Friedman and McCarthy (2020)	LS	Implications of AI for employment law and regulation	AI, audits, bias, discrimination, employment law, hiring, machine learning, regulation
Fritts and Cabrera (2021)	SS	Implications of AI for dehumanisation in hiring	AI, dehumanisation, employee-employer relationship, ethics, evaluation, human judgment, recruitment, screening
Fuchs (2023)	LS	Implications of AI for employment law and regulation	AI, anti-discrimination laws, bias, discrimination, fairness, hiring, legal accountability, New York City law, regulation
Galerpin (2019)	SS	Implications of online hiring platforms for gender discrimination	Gender segregation, gender stereotypes/discrimination, gig economy, online hiring
Gelles et al. (2018)	SS	Perceptions of AI	AI, Applicant Tracking Systems (ATS), complexity, fairness, hiring, perceptions, transparency, trust

Geyik et al. (2019)	CS	Measuring and mitigating algorithmic bias	AI, debiasing, demographic parity, equality of opportunity, fairness-aware ranking, quantification and measurement of bias, LinkedIn talent search, protected attributes
Glymour and Herington (2019)	CS	Measuring and mitigating algorithmic bias	AI, algorithmic bias, bias mitigation, behavior-relative error bias, disparate impact, disparate treatment, procedural bias, outcome bias, score-relative error bias
Gonen and Goldberg (2019)	CS	Measuring and mitigating bias in word embeddings	Debiasing, gender biases, natural language processing (NLP), text corpora, word embeddings
Hacker (2018)	LS	Implications of AI for employment law and regulation	AI, algorithmic audits, algorithmic decision-making, algorithmic discrimination, algorithmic fairness, data protection law, GDPR, EU anti-discrimination law, regulation
Hensler (2019)	LS	Implications of AI for employment law and regulation	AI, algorithmic discrimination, anti-discrimination law, auditing, data, discrimination, equality, GDPR, proactive regulation
Huang and Rust (2018)	HRMOS	Human-AI collaboration	AI, human-AI collaboration, innovation, mechanical/analytical/intuitive and empathetic intelligence, service, task
Huang et al. (2019)	HRMOS	Human-AI collaboration	AI intelligences: mechanical, thinking, and feeling, analytical, cognitive, 'feeling economy', emotional, empathetic, human workers, interpersonal, tasks
Jarrahi (2018)	HRMOS	Human-AI collaboration	AI, decision making, 'human-AI symbiosis', human augmentation, machine learning, organisational decision-making, processing capacity
Joyce et al. (2021)	SS	Implications of AI for inequality	AI, code, data, inequality, sociology of AI, structural social change
Kaminski (2018)	LS	Algorithmic governance	AI, algorithmic accountability, binary governance, decision-making, dignitary, justificatory, and instrumental concerns, GDPR, privacy, regulation
Kelan (2023)	HRMOS	Implications of AI for hiring	AI-supported hiring, 'algorithmic inclusion', bias, data, design, decisions, diversity, fairness, hiring, inequalities, machine learning, predictive algorithms
Kellogg et al. (2020)	HRMOS	Implications of AI for work	AI, algorithmic control, algorithmic occupations, organisational control, power, worker autonomy
Kim (2019)	LS	Implications of AI for discrimination and regulation	AI, anti-discrimination law, bias, big data, hiring algorithms, recruitment, protected groups, workplace
Kim (2020)	LS	Implications of AI for discrimination and regulation	AI, autonomy, bias, data, equality, fairness, hiring, inequality, labour markets, liability, online manipulation, opportunity markets, predictive algorithms, regulation, transparency
Köchling and Wehner (2020)	HRMOS	Algorithmic decision-making	Algorithmic decision-making, discrimination, fairness, HR development, HR recruitment

Kordzadeh and Ghasemaghahi (2022)	HRMOS	Implications of algorithmic decision-making for bias and discrimination	Algorithmic accountability, bias, data-driven decision making, discrimination, ethics, fairness, information systems
Kuhlman et al. (2020)	CS	Ethical implications of AI	Algorithmic bias, computing, data, diversity, ethics, fairness, representation, structural inequalities, systemic structural biases
Langenkamp et al. (2020)	CS	Fairness in algorithmic hiring	AI, algorithmic hiring, automation, fairness, employment decisions, hiring, legal and technological perspectives, machine learning, regulation, transparency
Lee et al. (2015)	CS	Algorithmic management	Algorithmic management, data-driven management, human workers, machines, transparency, work practices
Lee (2018)	SS	Perceptions of AI	AI, algorithmic management, decision making, emotion, fairness, human empathy, human skills, mechanical skills, perceptions, trust
Lepri et al. (2018)	CS	Algorithmic decision-making	Accountability, algorithmic decision-making, fairness, machine learning, technical solutions, transparency
Lin et al. (2020)	SS	Implications of AI for implicit bias	AI, engineering equity, human decision-making, human-machine interaction, human resource recruitment, implicit bias
Madaio et al. (2022)	CS	Perspectives of AI practitioners	AI systems, ethics, fairness, practitioners, organisational factors
Mann and Matzner (2019)	LS	Implications of algorithmic profiling for regulation	Algorithmic profiling, anti-discrimination law, bias, complexity, data protection, discrimination, invisibility, protection, regulation
Marti et al. (2024)	HRMOS	AI implementation within organisations	AI, disruptive algorithms, envelopes, fairness, organisational fields, regulation
Martin (2019)	HRMOS	Ethical implication of AI	AI accountability, algorithms, bias, decision-making process, ethics, responsibility, transparency
Michailidis (2018)	SS	Implications of AI for HR	AI, bias, blockchain, data, employment, HR, recruitment
Nachbar (2021)	LS	Algorithmic discrimination and fairness	AI, accountability, algorithmic fairness, computational considerations, decision-making, discrimination, legal considerations, regulation, transparency
Nawaz (2019)	HRMOS	Implications of AI for HR	AI, automation, bias, communication, data, decision-making, hiring, human bias, recruitment, screening
Newman et al. (2020)	HRMOS	Implications of AI for procedural fairness in hiring processes	Algorithmic decision-making, algorithmic reductionism, bias, fairness, HR decisions, human bias, procedural fairness/justice
Osoba et al. (2019)	SS	Implications of AI for equity and fairness	AI, algorithmic equity, algorithmic bias, decision-making, decision pipeline, fairness, hiring, recruitment, social applications, transparency
Parviainen (2022)	LS	Implications of AI for employment law and regulation	Algorithmic recruitment, automated decision-making, EU, GDPR, law, regulation
Raghavan and Kim (2023)	LS	Algorithmic discrimination and limitation of four-fifths rule	Algorithmic bias, algorithmic hiring, anti-discrimination law, discrimination, fairness, four-fifths rule

Raghavan et al. (2020)	CS	Algorithmic bias mitigation	Algorithmic bias, algorithmic hiring, anti-discrimination law, bias, de-biasing, disparate treatment/impact, technical perspectives, training data
Raisch and Krakowski (2021)	HRMOS	Human-AI collaboration	AI, automation-augmentation paradox; decision-making, human bias, human involvement, management, tasks, statistical bias
Rigotti and Fosch-Villaronga (2024)	LS	Implications of AI for fairness in recruitment	AI, anti-discrimination law, bias, data protection, discrimination, fairness, recruitment, regulation, transparency
Robert et al. (2020)	SS	Implications of AI for fairness in management	AI, autonomy, bias, decision-making, fairness, organisations, trust
Roemmich et al. (2023)	CS	AI as a technical solution for organisational problems	AI, bias, 'emotion artificial intelligence', hiring, organisational problems, technosolutions
Rosenblat et al. (2014)	SS	Implications of data-driven recruitment for discrimination	Algorithms, Applicant Tracking Systems (ATS), bias, data-driven recruitment, data, 'networked' employment discrimination
Rovatsos et al. (2019)	SS	Implications of algorithmic decision-making for bias and discrimination	Algorithmic bias, bias mitigation, data protection, decision-making, discrimination, fairness, GDPR
Shrestha et al. (2019)	HRMOS	Human-AI collaboration	AI, human-AI collaboration, organisational decision-making
Sonderling et al. (2022)	LS	Implications of AI for employment law and regulation	AI, accountability, anti-discrimination law, bias, employment discrimination, fairness, regulation, transparency
Tambe et al. (2019)	HRMOS	Implications of AI for HR	AI, bias, decision-making, discrimination, HR, performance evaluation, talent acquisition
Todoli-Signes (2019)	SS	Implications of AI for reinforcing bias/discrimination	AI, algorithms, automated decision-making, data protection, discrimination, GDPR, governance, regulation
Upadhyay and Khandelwal (2018)	HRMOS	Implications of AI for hiring	AI, automation, bias, hiring, recruitment, repetitive tasks
Wajcman (2017)	SS	Social and economic implications of technological change	AI, automation, future of work, politics of technology, power, robotics
Weiskopf and Hansen (2023)	HRMOS	Algorithmic governmentality	Accountability, algorithmic governmentality, bias, decision-making, ethics, people analytics, transparency
Xiang (2021)	LS	Legal and technical approaches to algorithmic bias	AI, algorithmic bias, bias mitigation, anti-discrimination law, data, fairness, legal perspective, liability, technical perspective
Yu et al. (2018)	CS	Ethical implications of AI	AI, accountability, bias, ethics
Yu (2020)	LS	Implications of AI for inequality	AI, accountability, 'algorithmic divide', bias, data, 'digital divide', equality, governance, regulation
Zuboff (2019)	SS	Social and economic implications of technological change	Behavioural traces, data extraction, digitisation, power, privacy, surveillance

¹ The concepts that appear in quotation marks are original terms used by the authors

Supplemental Materials C: Further Information on the Two-Phase Review Process

Introduction

In this supplemental material, we elaborate on the methods and decisions underpinning our two-phase review process: a scoping review (Phase 1) and a problematising review (Phase 2). This hybrid methodology combines the strengths of both approaches. The scoping review allowed us to establish our multidisciplinary review corpus in a systematic fashion, capturing the breadth of the emergent scholarship on AI, hiring and inequality. The problematising approach enabled us to go beyond mere description and engage more deeply with scholarship, to establish a critical multidisciplinary dialogue.

Our hybrid approach is particularly valuable because of the complexity of scholarship on AI, hiring and inequality (e.g., rapidly emerging, dynamic, multidisciplinary, multimethod). In the sections below, we offer a step-by-step justification and elaboration of our review approach. We begin by addressing the disciplinary positionality of our four-author team and our epistemological approach. We then provide details on the two phases of our hybrid approach.

Positionality

This review was conducted by a team of social scientists with overlapping expertise and shared interests in social, organisational and labour market inequalities. Our perspective has been shaped by collaboration with scholars from other disciplines—computing, statistics, management, and information systems—in a larger project examining the implications of AI for labour market (in)equality. The specific impetus for a multidisciplinary review emerged as a result of our experiences working on this larger project, including both challenges encountered

and solutions developed, to establish a common, situated, and systematic understanding of the role of AI in (re)producing labour market inequalities.

We follow the ideas of Alvesson and Sandberg (2020) that the reviewer is never neutral in guiding the review (p. 1295) and that reflexivity is central to knowledge production (p. 1297). Here we discuss how our positionality as social scientists shaped the way we interacted, analysed and understood the scholarship emerging across various disciplines on the role of AI in shaping/reshaping hiring processes and inequality embedded in these processes.

Our engagement with the topic of AI, hiring, and inequality is shaped by a critical lens that incorporates reflexivity and attention to structure. We view inequality as a structural, pervasive societal issue (e.g., Acker's inequality regimes) and focus on the role of technology in interacting with and reproducing it. From this perspective, AI systems are not neutral tools, but are shaped by—and help to reproduce—historical inequalities tied to gender, race, and class, gender. For example, we reflect on how the notion of 'algorithmic bias' is often framed in narrow or technical terms, without considering the deeper, pervasive social structure that underlies it.

As discussed in the paper, our analytical process in Phase 2 of the review (the problematising phase) involved developing and accumulating expertise within each disciplinary cluster by assigning and maintaining a 'primary reader' role, where the reader immerses themselves in the particular disciplinary cluster (e.g., CS, HRMOS, SS, and LS). Additionally, each author rotated across other disciplinary clusters as a 'secondary reader' to facilitate an interdisciplinary understanding of this scholarship.

While each of us has developed expertise on the CS, HRMOS, SS or LS scholarship for the duration of our review and as a part of our bigger project, we also acknowledge that our

original expertise, training and experience have shaped our evolving individual and collective understanding of the intra- and inter-disciplinary insights concerning AI, hiring and inequality. For example, our positionality is reflected in our interest in and focus on hiring as a foundational asymmetric process within market-based, capitalist economies that structures access to opportunities and has profound implications for the lives of individuals, their well-being, and economic security. Our choice of the key texts for problematising reflects our awareness of, and interest in, the pervasiveness of organisational inequality (Acker, 2006), within capitalistic structures marked by power asymmetries (Zuboff, 2019; Pasquale, 2015). Finally, our choice of hybrid methodology that combines scoping and problematising methods highlights our interest in *how* knowledge is being created in addition to *what* knowledge is being created (Alvesson and Sköldbberg, 2018)

Epistemology

Our epistemological stance draws on ideas of reflexive methodology (Alvesson and Sköldbberg, 2018), which emphasises the importance of moving between empirical analysis, theoretical interpretation, and reflection on our own role as researchers. Following our disciplinary positionality as scholars of social, organisational, and labour market inequality, our review is inspired by critical and constructivist understandings of knowledge production. We recognise that knowledge about AI, hiring, and inequality is not simply discovered but is actively constructed through debates, across disciplines, domains and stakeholders, and that different approaches and assumptions guide these debates. Thus, our review is not a neutral summary or description of a found field of knowledge, but is a situated contribution embedded in our disciplinary backgrounds and research decisions. Consequently, our choices—what we include,

what we leave out and how we analyse data—are shaped by our standpoints, as scholars and individuals. As social scientists, we approach AI, hiring, and inequality not as isolated technical phenomena, but as part of broader organisational, social and political structures. As individuals residing in the Global North, with diverse transnational origins, educational training, and experiences, we are aware of how our perspectives are shaped by the debates and discourses that prevail in this part of the world. Our status as academics facilitates participation in these debates.

Alvesson and Sandberg (2020) underscore the reviewer’s role “as an artist, a detective, an innovator or even an anthropologist, supporting the innovative part of research” (1302). In line with their problematising methodology, we go beyond ‘representing’ a field of scholarship. Using a hybrid approach we aim to utilise scoping techniques to map the emerging knowledge and its construction, and to utilise problematising techniques to critically engage with the assumptions that underlie how AI and inequality are currently studied in the context of hiring and organisations. Our goal is thus to ‘open up’ space for new questions and to shift how key issues are being approached in this field.

Overview of our hybrid process

Phase 1: Scoping review:

The goal of this phase was to identify the breadth of emerging ideas and debates on the intersecting domains broadly defined as AI, hiring practices and inequality (AI/hiring/inequality). A scoping methodology was chosen as the most suitable method for assembling literature and mapping key ideas in rapidly emerging fields, which are not stable and complex (Mays et al., 2001:194). The knowledge emerging at the intersection of AI, hiring

practices, and inequality is such a field, and a scoping approach was essential to assess the breadth of the field to paint a representative picture of the emerging knowledge. The scoping approach combines rigour with flexibility, allowing for the collection of emerging work across disciplines and methods (Arksey and O'Malley, 2005; Daudt et al., 2013). To maintain rigour in the process of selecting and mapping the scholarship across disciplines and methods, we adapted the 5-stage framework of Arksey and O'Malley (2005), who contributed to the methodology of the scoping review and identified ways in which the scoping method is different and similar to the systematic review process.

Step 1: Formulating research questions: To begin, we formulated several broad questions that our review seeks to answer: What disciplines are involved in studying the role of AI in hiring and recruitment? What questions are being asked? What is known about the capacity of AI to reproduce, mitigate or form new forms of inequality? These questions guide our subsequent literature search and review strategies. Following scoping methodology, our questions were broad enough to capture the breadth of the emerging ideas and debates.

Step 2: Search: To operationalise our questions, we used the combination of keywords from the following areas: 1) *algorithms* and *AI*; 2) *human resource management, hiring, recruitment, work, employment* and *organisational* processes; and 3) *(in)equality, fairness, bias, stereotypes, prejudice* and *discrimination* (we will refer to it as an intersection of 3 domains: AI/hiring/inequality) to search for literature on Google Scholar, followed by ProQuest, Sociological Abstracts, Web of Science Social Sciences Citation Index, Sage Journals, and Social Sciences Research Network.

We used broad criteria to identify distinct types of publications, questions, and ideas, and we did not incorporate filters of ‘quality’ that are typical for systematic reviews (e.g., limiting the corpus to certain databases). This approach is particularly suitable for our multidisciplinary review due to different publication traditions across different disciplines. For example, CS papers are typically published as peer-reviewed conference proceedings indexed by Google Scholar but not by conventional databases such as Scopus and the Web of Science. In this case, a database-informed approach, though widely used in previous single-discipline reviews, would not be suitable for establishing a quality filter for our corpus. Rather, we more specifically ensured that the individual entries included in our review corpus are of high quality— e.g., they are published in widely recognised and credible outlets (e.g., journals, conference proceedings). We screened titles and abstracts to identify the substantive relevance of the items to the subject matter. This search resulted in 389 publications.

Step 3: Selection and screening: The 389 publications were exported to the Rayyan platform to allow for blind inter-rater screening. Our authorship team then screened these 389 publications, using individual blind selection. Each author screened the 389 publications and assigned the label of either ‘included’ (criteria: highly relevant to our review), ‘excluded’ (criteria: not very relevant to our review) or ‘maybe’ (criteria: somewhat relevant to our review). The notion of ‘relevance’ was defined as substantive relevance to the intersection of 3 domains of interest (AI/hiring/inequality). The ‘relevance’ was defined as a substantive engagement with the 3 domains of interest. For example, publications that discussed AI and inequality in the context of hiring/workplaces/organisations were tagged as ‘highly relevant’.

Inclusion criteria: inter-rater agreement on the highest relevance to the intersection of the domains, broadly defined as AI/hiring/inequality.

Exclusion criteria: inter-rater agreement on low-medium relevance to the intersection of the domains, broadly defined as AI/hiring/inequality.

The final decision on a label was based on the majority coding. In the case of a tie between labels, we adopted an inclusive approach to ensure that we did not miss out on important publications. For example, if the tie was between ‘excluded’ and ‘maybe’, the final decision would be ‘maybe’. If the split was between ‘maybe’ and ‘included’, the final decision was ‘included’.

Of the items reviewed, 75% reached a majority agreement (3/4 reviewers) reflecting strong inter-rater reliability in the selection process. We proceeded with the category of included, which consisted of 212 empirical, conceptual, and applied publications in different disciplines.

Step 4: Charting the data: We then created a detailed mapping of the 212 publications, noting discipline, key questions, concepts and insights. Based on this mapping, we identified four disciplinary clusters connected to 1) computer sciences (CS); 2) human resource, management and organisation studies (HRMOS); 3) other social sciences (SS); and 4) legal scholarship (LS). As noted in our paper, we do not view the borders of disciplinary clusters as rigid, given the complexity and novelty of emerging knowledge and given the emerging communication across disciplines.

The cluster that we identified as **Computing Science (CS)** is limited to literature that incorporates social aspects/concerns/perspectives along with technical perspectives, thus excluding a vast computer science literature that engages only with technical perspectives. The

computer science discipline currently dominates knowledge generation on AI, focusing predominantly on the technical aspects of algorithms. Therefore, it is important to note that our review engages with the slice of this scholarship that speaks to social perspectives along with technosolutions.

The cluster that we identify as **Human Resource, Management, and Organisation Studies (HRMOS)** contains both applied and critical writings that centre on human resource practices in organisations, with a special focus on AI applications and implementation in HR processes. We identify this scholarship as a unique cluster and single it out from the broader social sciences (SS) debates.

The cluster we identify as **Social Sciences (SS)** is a heterogeneous cluster that combines discussions and debates from various fields of sociology, psychology, philosophy, and science and technology studies (STS). While coming from diverse disciplines, these debates offer a broader, critical lens on matters of AI and the social implications of its usage.

The cluster we identify as **Legal Scholarship (LS)** offers very focused writing on legal issues, centring on Western, North American, and European perspectives. The regional foci of this body of literature are partly a result of our focus on publications in the English language only, which we have acknowledged as a limitation of our review in the main article.

In a small number of cases where an item could potentially belong to more than one cluster (i.e., socio-technical perspectives which draw expertise from the CS and SS), we assigned it to the cluster that was most substantively relevant, such as the key focus of the paper, the scholarship that the paper engages with and the expertise of the authors.

The distribution of the 212 items by discipline is as follows:

Discipline	Number	Percentage
CS	47	22.2%
HRMOS	73	34.4%
SS	55	25.9%
LS	37	17.5%
Total	212	100.0%

Phase 2: Problematising review

Phase 1 highlighted that the literature emerging on the intersection of AI/hiring/inequality is multidisciplinary and multi-method, and it touches on the phenomena (AI/hiring/inequality) from various angles, fragmented and therefore very complex. The goal of Phase 2 was to gain a deeper understanding of these debates and to critically interrogate the underlying assumptions that drive questions and concerns across the disciplinary clusters we have identified. The goal was to engage with the scholarship more analytically, selectively, and critically, problematising *what* is known and *how* it is known, with the attempt to identify novel ways of understanding and conceptualising the role of AI in hiring and inequality.

In this phase, we followed a problematising review, developed by Alvesson and Sandberg (2020), which was chosen as the most suitable methodology for this phase, as it provides a methodological and epistemological framework for the deep, reflexive analysis of the scholarship. This type of analysis goes beyond a scoping review, which focuses on breadth rather than depth. As Alvesson and Sandberg (2011: 32) clarify, the problematising is not ‘an end in itself’ but rather ‘means to identify and challenge assumptions underlying existing theory and, based on that, being able to formulate more informed and novel research questions’.

In this phase, we followed four core principles of problematising review: 1) ‘reflexive reading’; 2) ‘less is more’; 3) ‘reading broadly but selectively’; and 4) ‘not accumulating but problematising’ (Alvesson and Sandberg, 2020, p. 1297).

Principle 1: Alvesson and Sandberg (2020) suggest that “reflexivity typically calls for the researcher to read a limited number of texts carefully, to challenge his or her interpretations by considering alternative perspectives and sources of inspiration, to work with doubt and recognise intuition, and to aim for insightfulness rather than rigour or pseudo-rigour” (p. 1297).

Recognising the importance of reflexivity, in Phase 2, we relied more heavily on the knowledge acquired in the first stage of the review and, deep, reflexive and layered reading in the second phase of the review.

We practised reflexivity by assigning a ‘primary reader’ to each disciplinary cluster. The ‘primary reader’ immersed themselves in the disciplinary cluster to develop and maintain the intra-disciplinary ‘expertise’ throughout the duration of the selection and analysis conducted in phase 2. Additionally, we assigned ‘secondary readers’ who rotated across disciplinary clusters, providing further insight. This exercise allowed us to generate and accumulate insight within and across the disciplinary clusters and identify connections between them, the similarities and differences of the ideas, the underlying assumptions that shape these ideas, and the complex ways in which these ideas and assumptions interact. As in the case of reflexive, critical qualitative data analysis and coding, the primary and secondary readers first assessed their assigned publications separately. They then engaged in deliberation to finalise their analysis of the supercorpus. This deliberation, sustained through numerous written and oral discussions, allowed the reviewers to reflect on their disciplinary/expertise positionality, perspectives and

subjectivity. For example, our finding on the asymmetry in knowledge generation and the domination of the computer science perspective was only possible to be identified via multiple layered readings and interpretations of knowledge generated inside and across clusters.

Principle 2: Following the idea of ‘less is more’, we created a ‘supercorpus’ – a selection of readings from the scoping review stage (57 items) and additional readings from reading beyond the corpus (40 items).

The 57 items selected from the original corpus represent key ideas and debates emerging from each discipline. The key criteria for inclusion/exclusion in Phase 2 were identifying the most insightful contribution, such as representing a key debate, initiating a new concept/debate, and featuring high substantive relevance to the intersection of AI/hiring/inequality.

As a practice, this selection was operationalised in accordance with the characteristics of the cluster itself. We relied on the accumulated expertise of the ‘primary reader’ – the reviewer who immersed themselves in the original disciplinary cluster and gained a broad understanding of the key questions and debates. Each ‘primary reader’ suggested a list of 10-15 articles from their disciplinary cluster for deeper examination. It is important to note, that, epistemologically, our focus in this phase shifted towards a deep, reflexive, layered reading of the small collection of items, in such a process, as previously mentioned, the reviewer is not seen as a neutral ‘puzzle solver’ but rather a ‘creative artist’ whose previous and emerging expertise and positionality shape the way they interact with the scholarship. To enhance rigour, each reviewer was rotating across other disciplinary clusters as a ‘secondary reader’, to meaningfully engage in interdisciplinary analysis.

The disciplinary breakdown of these 57 items is:

Discipline	Number	Percentage
CS	12	21.1%
HRMOS	14	24.6%
SS	15	26.3%
LS	16	28.1%
Total	57	100.0%

Principle 3: Following the idea of reading broadly but selectively, we also read beyond the initial corpus and searched for additional scholarship in each discipline, with the aim of expanding and updating the identified debates.

This strategy is adopted for several reasons: (1) AI and inequality in labour market processes is a rapidly evolving field; as we worked on our review paper, new important publications emerged; (2) ongoing research on AI also relates to classic literature that is not directly related to AI but concerns technological change and its implications for organisational inequalities; (3) our core selection was strongly tied to the intersection of 3 domains (AI/hiring/inequality), while other items, in neighbouring or overlapping domains, could be relevant for a deeper understanding of current debates. Our additional search process allowed us to incorporate these key texts and to bring them to bear on the emerging scholarship on AI.

Thus, reading ‘beyond corpus’ was operationalised via 1) searching the most recent publications related to the intersection of AI/hiring/inequality; and 2) hand searching for items that will help to deepen the discussion on some areas of the intersection, such as broader sociological discussion on implications of AI for society and social inequality that are not

directly relevant to hiring. Part of this process involved conducting an additional sweeping search to identify debates that are more conceptual, theoretical and innovative. Similarly, this process of selection was driven by the notion of reflexivity and sustained by the reviewers with accumulated knowledge and expertise across and within disciplinary clusters. As noted previously, epistemologically, ‘problematism method’ challenges ideals such as rationality, procedure, transparency, and being trustful of conventions” (Alvesson and Sandberg, 2020, p. 14), instead, it centres on the author’s expertise, creativity and reflexivity in the process of selecting and analysing data. We collected an additional 40 items to be added to our ‘super corpus’ for in-depth examination.

The disciplinary breakdown of these items is:

Discipline	Number	Percentage
CS	7	17.5%
HRMOS	18	45.0%
SS	5	12.5%
LS	10	25.0%
Total	40	100.0%

The total number of items forming a ‘supercorpus’ and reviewed in Phase 2 was 97, as reported in the table below, which consists of the selection of readings from the scoping stage and from ‘reading beyond corpus’. The breakdown of the total ‘supercorpus’ is as follows:

Discipline	Number	Percentage
CS	19	19.6%

HRMOS	32	33.0%
SS	20	20.6%
LS	26	26.8%
Total	97	100.0%

Principle 4: While in Phase 1, our key goal was to assess the breadth of the accumulating knowledge, in Phase 2, we took an approach of problematising, which Alvesson and Sandberg (2020) define as an “‘opening up exercise’ that enables researchers to imagine how to rethink existing literature in ways that generate new and ‘better’ ways of thinking about specific phenomena” (page 1291) with the goal of re-evaluating (instead of integrating) existing understandings of phenomena (Alvesson and Sandberg, 2020, p. 1295). This process is identified by analytical and creative rather than systematic evaluation of ideas, and it aims to question and search beneath the more salient ideas.

To facilitate this process, we engaged with theoretical frameworks that help to provide a conceptual lens to understanding the (pre-AI) pervasiveness of inequality in hiring and organisations (Acker, 1990) and the capacity of AI to conceal and reproduce power asymmetries (Zuboff, 2019; Pasquale, 2015). In this step, we used these foundational readings as a prism to gain a deeper understanding, but also ‘develop an alternative assumption ground with the potential to become the start of a novel theoretical contribution’ (Alvesson and Sandberg, 2020, p. 1300) of how AI interacts with inequality regimes.

A central ambition of problematising is ‘to generate re-conceptualizations of existing thinking that trigger new ideas and theories’ (Alvesson and Sandberg, 2020, p. 1295). Our review and our engagement with the conceptualisations of Acker (1990), Zuboff (2019), Pasquale (2015) and others allowed us to introduce a new concept – ‘algorithmically-mediated’

inequality regimes to capture the way AI interacts and embeds itself within the inequality regimes conceptualised by Acker (1990) which are now cementing, concealing and legitimising the inequality mechanisms within the algorithmic complexity and pervasiveness. Our problematising exercise demonstrates that the ‘algorithmically-mediated’ regimes are identified by increased ‘algorithmic invisibility’ and almost ubiquitous pervasiveness, and, therefore, legitimacy, of algorithmic solutions.

The table below provides a summary of articles by discipline at each phase of the review process.

Hybrid Review: Distribution of articles by cluster at each stage of review						
PHASE 1	Phase 1 Scoping review (articles identified)					
		CS	HRMOS	SS	LS	Total
	Number	47	73	55	37	212
	% of Total	22%	34%	26%	17%	100%
PHASE 2	Phase 2a Problematising review (articles selected from initial scoping review above)					
		CS	HRMOS	SS	LS	Total
	Number	12	14	15	16	57
	% of Total	21%	25%	26%	28%	100%
	Phase 2b Problematising review (articles identified by reading beyond the corpus)					
		CS	HRMOS	SS	LS	Total
	Number	7	18	5	10	40
	% of Total	18%	45%	13%	25%	100%

FINAL	Final Supercorpus (total articles identified)					
		CS	HRMOS	SS	LS	Total
	Number	19	32	20	26	97
	% of Total	20%	33%	21%	27%	100%

References

- Acker J (1990) Hierarchies, jobs, bodies: A theory of gendered organizations. *Gender & Society* 4(2): 139-158.
- Acker J (2006) Inequality regimes: Gender, class, and race in organizations. *Gender & Society* 20(4): 441-464.
- Alvesson M and Sandberg J (2020) The problematizing review: A counterpoint to Elsbach and Van Knippenberg's argument for integrative reviews. *Journal of Management Studies* 57(6): 1290-1304.
- Alvesson M and Skoldberg K (2018) *Reflexive Methodology: New Vistas for Qualitative Research*. London: Sage Publications.
- Arksey H and O'Malley L (2005) Scoping studies: Towards a methodological framework. *International Journal of Social Research Methodology* 8(1): 19-32.
- Daudt HML, van Mossel C and Scott SJ (2013) Enhancing the scoping study methodology: A large, inter-professional team's experience with Arksey and O'Malley's framework. *BMC Medical Research Methodology* 13(48): 1-9.
- Mays N, Roberts E and Popay J (2001) Synthesising research evidence. In: Allen P, Black N, Clarke A, Fulop N, Anderson S (eds) *Studying the organisation and delivery of health services*. London: Routledge, 188-220.

- Pasquale F (2015) *The Black Box Society: The Secret Algorithms That Control Money and Information*. Cambridge, MA: Harvard University Press.
- Sandberg J and Alvesson M (2011) Ways of constructing research questions: Gap-spotting or problematization? *Organization* 18(1): 23-44.
- Zuboff S (2019) *The Age of Surveillance Capitalism: The Fight for a Human Future at the New Frontier of Power*. New York, NY: Public Affairs.