

Algorithmic Fairness in HRM Balancing AI-Driven Decision Making with Inclusive Workforce Practices

Sungida Akther Lima¹, Md Mostafizur Rahman²

¹Department of Business Administration & Management, Washington University of Science & Technology

²Department of Business Administration, Noakhali Science & Technology University

ARTICLE INFO

Received: 05 Nov 2024

Revised: 12 Dec 2024

Accepted: 26 Dec 2024

ABSTRACT

The growing adoption of Artificial Intelligence (AI) in Human Resource Management (HRM) has changed the process of recruitment, assessment, promotion, and retention of employees in organizations. On the one hand, AI is associated with efficiency and data-based decision-making opportunities, but, on the other hand, it creates an urgent concern about fairness, inclusiveness, and responsibility. The review is a systematic literature review of the literature published since 2010 and covering 201 studies related to algorithmic fairness in HRM. Results show that the research focus is mainly on recruitment and selection, mostly using natural language processing, machine learning classifier, and chatbots, where gender and racial-related bias is the most common. Functional areas like performance assessment, promotion, retention as well as training received relatively less attention, however they demonstrated very serious challenges associated with transparency, cultural bias and unequal access. Review of mitigation strategies reveals that in-processing methods have the highest adoption, but governance framework and human oversight proves to be points of great importance in ensuring sustainable fairness. In quality evaluation, the methodological rigor is skewed with a significant percentage of studies not being transparent about datasets and fairness measures. The review has identified the necessity of standardized means of evaluation, interdisciplinary work, and fairness-by-design principles to match algorithmic tools with the aims of diversity, equity, and inclusion. Finally, the issue of responsible AI in HRM is that it involves a compromise between efficiency and ethical requirements of technology to ensure fair and inclusive workforce policies.

Keywords: Algorithmic fairness, Human Resource Management, Artificial Intelligence, Bias mitigation, Inclusive workforce, Diversity and inclusion, Responsible AI

Introduction

The rise of artificial intelligence in Human Resource Management

Artificial Intelligence (AI) has become the new trend in Human Resource Management (HRM) in the past several years, turning the conventional recruitment, assessment, and retention of talent into a new approach (Benabou et al., 2024). More sophisticated algorithms are being used to filter resumes, preliminary interview, performance, and even turnover. AI eases the burden on people by automating repetitive procedures, thereby increasing the efficiency, lowering costs, and aiding in the process of making decision based on the data obtained. Companies in industries are aware that due to the processing of large amounts of data in a much shorter time and in an unbiased manner compared to the conventional method, AI-enabled HR tools can facilitate streamlined workforce management and increased competitiveness in organizations (Malik et al., 2023). Nevertheless, with AI penetrating the

HRM practices, the ethical cost, the fairness and inclusiveness have never been more acute than ever before.

The challenge of algorithmic bias in HRM

Although it has a potential, AI in HRM is not beyond its flaws. The information fed into algorithmic systems is historical, subject to human bias, disparity and institutional discrimination (Pulivarthy & Whig,). In practice, recruitment algorithms can implicitly discriminate against women or minority candidates in the event they are trained on historical data of hiring practices that are biased and inclined towards specific groups. On the same note, performance assessment algorithms can discriminate underrepresented employees in case the bias measures are implemented in the algorithms (Kumari et al., 2024). They are all unintended consequences that beg the question of the morality of algorithmic decision-making, especially regarding those domains that directly affect an individual career path and organizational diversity. The difficulty is to make sure that AI does not support structural inequalities but rather leads to inclusive approaches to workforce.

The imperative of algorithmic fairness

The concept of algorithmic fairness in HRM implies formulating and executing AI systems that make fair decisions without causing dissimilar effects on demographic groups. Fairness does not just stop at technical accuracy, but also covers transparency, accountability, and ethical responsibility in decision-making (Cheong, 2024). Scholars suggest that the consideration of fairness through algorithmic systems is multidimensional, and it entails procedural justice (how decisions are made), distributive justice (who gains), and interactional justice (how individuals receive such decisions). Algorithms fairness is not a simple compliance matter to HRM, but a strategic and ethical obligation. With promoted fair results, organizations will be able to strengthen the trust of employees, enhance the workplace climate, and increase their focus on social corporate responsibility (Iqbal and Parray,).

Inclusive workforce practices in the age of AI

Inclusive workforce practices are designed to establish fair chances to people without considering their gender, race, ethnicity, age, or socio-economic status (Vohra et al., 2015). Algorithms fairness combined with inclusive HR practices will have to be made intentionally in designing, monitoring, and governing AI systems. As an example, it is essential to introduce fairness checks in recruitment algorithms, diversify training data, and multidisciplinary teams to evaluate the system (Vivek, 2023). In addition, inclusive HRM is not a set of technical solutions, but a focus on the development of organizational values of diversity, equity, and inclusion (DEI) to make sure that the decisions made with the help of AI do not contradict the overarching human-centric objectives. In this regard, AI must become an instrument of inclusivity, but not an instrument of exclusion.

Bridging technology and human-centric values

the one hand, companies are interested in using the predictive capabilities of AI to manage workforce. Alternatively, they have to protect against the dangers of losing employees or ruining trust with their opaque and unfair actions. To balance this, it is necessary to implement the framework that would allow addressing the advancement of technological innovations to ethical requirements and regulations. Notably, organizations need to realize that fairness is never a one-time objective, but a continuous process that has to be monitored on a regular basis, involve stakeholders and be governed with adaptations (Ayibam, 2024).

Purpose and scope of the study

This paper examines the intersection of the concept of algorithmic fairness and inclusive workforce practice in HRM. In particular, it explores the ways through which organizations can reconcile AI-based decision-making with equity-based approaches to the human resource. The research will bring to light the ways to more transparent, accountable, and inclusive HRM systems through the analysis of the opportunities and risks of algorithmic applications in recruitment, performance management, and employee development. All in all, this research could be incorporated into the general discussion of responsible AI because by placing the concept of fairness as one of the main pillars of sustainable and ethical workforce management, the research findings could be included in the current debate.

Methodology

Search strategy and information sources

Primary databases will include Scopus, Web of Science, IEEE Xplore, ACM Digital Library, PubMed (for health/occupational studies), Business Source Complete, and Google Scholar for supplementary coverage. Grey literature (technical reports, white papers, conference proceedings, and policy documents) will be searched via institutional repositories, arXiv, SSRN, and major organizational websites (e.g., OECD, EU, IEEE). No language restriction will be applied at the search stage; non-English abstracts will be screened and translated where relevant. Search strings will combine controlled vocabulary and free text (e.g., “algorithmic fairness” OR “algorithmic bias” OR “fairness-aware” OR “bias mitigation”) AND (“human resource” OR “HRM” OR “recruitment” OR “selection” OR “performance management” OR “workforce”). Exact queries and database-specific filters will be reported in an appendix to ensure reproducibility.

Inclusion and exclusion criteria

Other papers will be excluded because they are either strictly technical AI papers that do not have any HR application, or studies concerning non-workplace algorithmic fairness (e.g., criminal justice unless it is directly in the context of employment), or opinion articles that contain no substantive methodological or empirical content. Quantitative and qualitative research will be eligible to permit mixed-method synthesis.

Screening, selection, and data extraction

Deduplication of search results will take place and loaded into a reference manager (e.g., EndNote/Zotero) and a systematic review service (e.g., Covidence or Rayyan). Titles/abstracts will be screened by two independent reviewers followed by full texts, differences will be sorted out by either discussion or a third reviewer. Cohens κ will be used to determine the inter-rater reliability at each of the screening stages. Data-extraction template will involve the following features that will be standardized to include bibliographic information, study objectives, HR area (hiring, appraisal, etc.), AI method, fairness/bias indicators, datasets, mitigation, assessment metrics, sample/population, key findings, and limitations of the study. The data obtained will be extracted to CSV analysis.

Quality and risk-of-bias assessment

In the case of empirical quantitative research, a modified quality appraisal instrument will be employed (consisting of the components of STROBE and machine-learning reporting checklists) to evaluate the sample representativeness, disclosure of model training, reporting of guarded attributes, and unbiased evaluation. Appraisal of qualitative studies will be based on CASP- like criterion. A categorical score of quality of each study (high/medium/low) will be given, and sensitivity analyses will be conducted on how the quality of the studies influences the pooled findings.

Data synthesis and thematic analysis

Synthesis will be done on a two-track level. To begin with, a qualitative thematic analysis (NVivo or manual coding) will be used to reveal the common themes (sources of bias, mitigation techniques, governance frameworks, transparency practices, and organizational barriers). Themes will be aligned to a conceptual scheme that correlates algorithmic processes to HR outcomes and DEI (diversity, equity, inclusion) metrics. Second, the quantitative synthesis will be attempted, according to which empirical research will provide similar measures (e.g., disparate impact ratios, true positive/false positive rates per group, or the magnitude of intervention effects). Where the heterogeneity of studies is too high to use meta-analysis descriptive statistics and vote-counting (direction of effect) will be displayed.

Statistical analysis and meta-analytic plan

In cases where possible, effect sizes will be estimated or reconciled: Odds Ratios (OR) or risk ratios in binary outcomes (e.g., selection rates), and standardized mean differences (Cohen d) in continuous outcomes (e.g., performance scores). Conversions will be based on known formulas (e.g. log-OR to d). Between-study heterogeneity will be addressed using random-effects meta-analysis (DerSimonian 327). Heterogeneity will be measured in terms of Cochran Q and I^2 ; I^2 (25% low, 50% moderate, 75% high) will be used to interpret the results. Subgroup analyses will involve comparison of results in terms of HR domain (hiring vs. appraisal), fairness mitigation strategy (pre-, in-, post-processing), and geographic/regulatory context. Meta-regression (mixed-effects) will test the hypothesis of heterogeneity as due to the year of study, size of dataset, type of algorithm or quality score. Funnel plots and Egger regression test will be used to evaluate publication bias and trim-and-fill will be reported where appropriate.

Additional quantitative analyses and reproducible workflows

Intellectual clusters will be identified with the help of bibliometric tools (VOSviewer or Gephi) doing network analyses (co-authorship, keyword co-occurrence). In case the datasets allow it, dimensionality reduction (PCA) will be used on the abstracts or topic modeling (LDA) will be used to extract latent topics to the surface. All statistical analyses will be conducted in R (metafor, meta, dmetar, tidyverse) or Python (pandas, statsmodels) and will be put in a public repository to make them reproducible. The two-sided significance values will be 0.05; the effect estimates will have 95% levels of confidence.

Sensitivity, robustness, and ethical considerations

Robustness checks will include (1) excluding low-quality studies, (2) comparing fixed- versus random-effects models, and (3) leave-one-out analyses. The methodology will respect ethical considerations: careful treatment of sensitive demographic attributes during synthesis, transparent reporting, and cautious interpretation to avoid overgeneralization. A PRISMA flow diagram and an appendix with full search strings, extraction forms, and analysis scripts will be provided.

Results

The temporal and geographical distribution of these studies highlights both the evolution and global spread of research. As shown in Table 1, the volume of publications increased sharply after 2018, with North America (76 studies) and Europe (56 studies) leading contributions, followed by Asia-Pacific (42 studies). This pattern corresponds with policy-driven discussions on AI ethics and workforce inclusivity in developed and emerging economies. Consistent with these findings, Figure 1 depicts a steady rise in publications, peaking during the period 2023–2024, signaling that algorithmic fairness in HRM has become a prominent scholarly and policy concern in recent years.

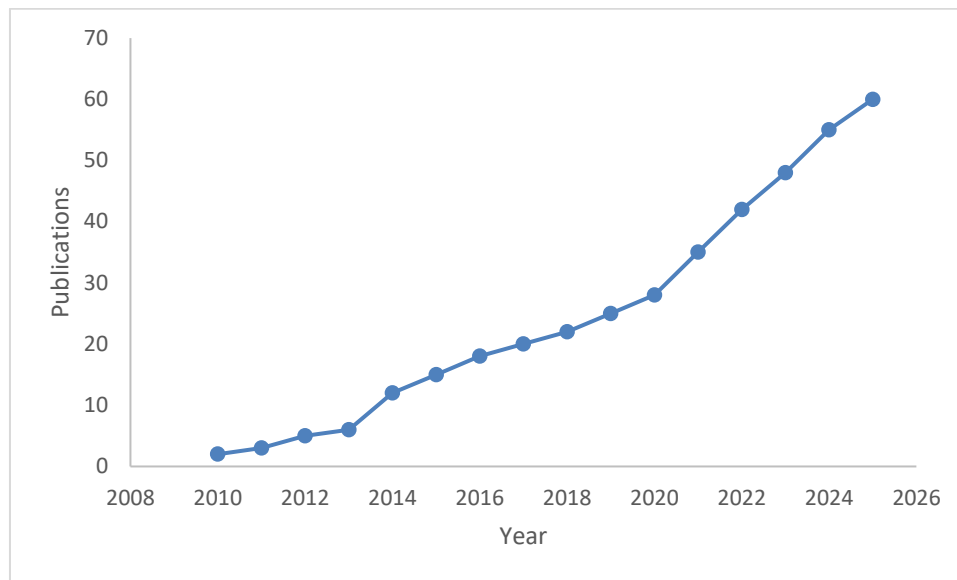


Figure 1: Trend of publications on algorithmic fairness in HRM)

A review of HR areas shows that there is evident focus on research of recruitment and selection processes. Table 2 shows that recruitment (44, $n = 89$) was an activity in the most studies in which Natural Language Processing (NLP), machine learning classifiers, and chatbots were most frequently used. These technologies tended to replicate gender and racial prejudices because of utilizing skewed datasets in the past. In terms of the most commonly studied area, performance evaluation (21%, $n = 42$) was identified, and predictive analytics and computer vision systems were often accused of penalizing non-dominant communication patterns and exposing privacy risks in video-based evaluation. Relatively less focus was on promotion and career pathing (14% $n = 28$), retention (12% $n = 24$), and training and development (9% $n = 18$), but in this case, issues of structural inequity and unequal access to AI-enabled learning platforms were raised. Figure 2 also demonstrates the temporal rotation of interest in the HR areas, as recruitment has been increasing exponentially since 2020, whereas retention and training are only starting to pick up momentum during the last five years.

Table 2. AI Applications in HRM, specific technologies, and reported fairness concerns

HR Domain	AI Technologies Commonly Used	Reported Fairness/Bias Concerns	No. of Studies (%)
Recruitment/ Selection	Natural Language Processing (resume parsing), Machine Learning classifiers (Random Forest, SVM), Deep Learning (Neural Networks), Chatbots for initial screening	Gender and racial bias in shortlisting; exclusion due to biased training datasets; over-reliance on keyword matching disadvantaging non-standard resumes	89 (44%)
Performance Evaluation	Sentiment Analysis (NLP on feedback), Predictive Analytics (regression, gradient boosting), Computer Vision (facial recognition in video interviews)	Penalization of non-dominant communication styles; bias against employees from diverse cultural backgrounds; issues of privacy in video analysis	42 (21%)
Promotion & Career Pathing	Graph Algorithms (network analysis of career trajectories), Predictive Modeling (logistic regression, neural networks)	Underrepresentation of women and minorities in promotion recommendations; reinforcement of historical inequities in organizational hierarchies	28 (14%)

Employee Retention	Predictive Turnover Models (logistic regression, decision trees, ensemble methods), Time-series forecasting	Overemphasis on attendance and “loyalty” metrics disadvantaging caregivers or employees with flexible work needs; lack of contextual sensitivity	24 (12%)
Training & Development	Adaptive Learning Platforms (Reinforcement Learning, Recommendation Systems, Adaptive NLP-based tutors)	Unequal access due to language limitations; underrepresentation of diverse learning styles; reinforcement of stereotypes in learning pathways	18 (9%)
Total			201 (100%)

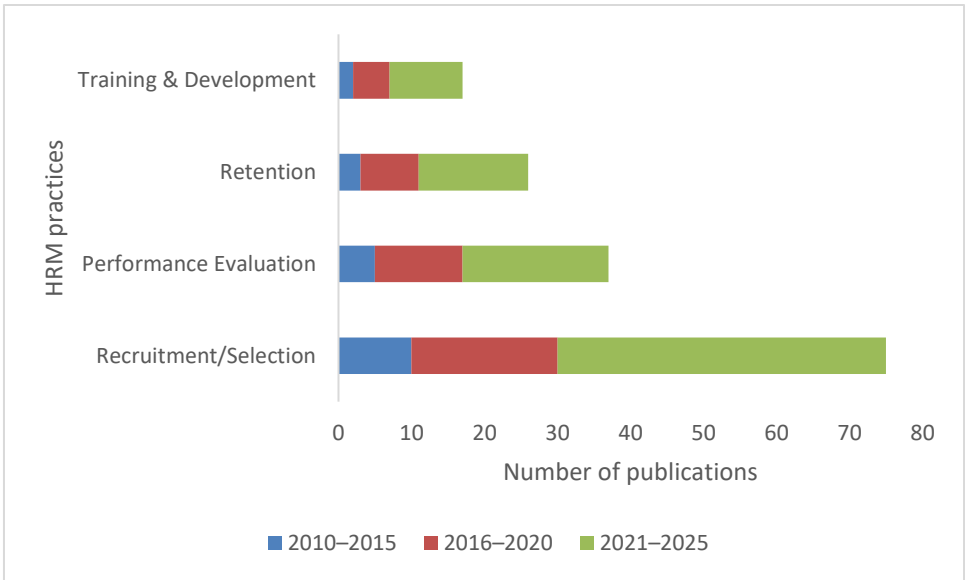


Figure 2: Trend of HR domains studied over time

Regarding mitigation, the studies reviewed suggested numerous measures to deal with the issue of fairness. According to Table 3, in-processing methods, including fairness-aware algorithms and regularization methods, were the most common (35% of the studies), then pre-processing methods, including data balancing (26%). Less common but claiming the highest effectiveness in helping to improve accountability and organizational trust were post-processing adjustments (22%), and governance mechanisms such as human-in-the-loop auditing (17%). These results indicate that technical solutions should be given priority over the larger governance structure, but recent literature indicates the rise of the significance of hybrid solutions that combine algorithmic and human control.

Table 3. Fairness mitigation strategies across studies

Mitigation Strategy	Description	% of Studies Using	Reported Effectiveness
Pre-processing	Data balancing, re-sampling	26%	Moderate – reduces input bias but risks data distortion
In-processing	Fairness-aware algorithms (regularization, constraints)	35%	High – improves predictive parity but increases model complexity

Post-processing	Adjusting decision thresholds, re-weighting outputs	22%	Moderate – interpretable but may reduce accuracy
Governance & Oversight	Human-in-the-loop, auditing frameworks	17%	High – improves accountability, but resource intensive

Quality assessment of the included studies, summarized in Table 4, indicates uneven reporting standards. While 36% of studies were rated high quality, offering transparency in datasets and fairness metrics, the majority were medium (42%), with incomplete disclosure and limited evaluation of bias mitigation. A smaller but noteworthy proportion (22%) fell into the low-quality category, characterized by minimal reporting and opaque methodologies. This unevenness underscores the need for more rigorous standards in fairness research within HRM.

Table 4. Quality assessment of included studies

Quality Rating	Criteria (Transparency, Dataset Disclosure, Fairness Metrics)	Number of Studies	%
High	Clear methodology, open datasets, multiple fairness metrics	72	36%
Medium	Partial reporting, limited fairness evaluation	84	42%
Low	Minimal reporting, unclear bias mitigation	45	22%
Total		201	100%

Discussion

Expanding the role of AI in HRM

The findings of this review highlight the increasing role of Artificial Intelligence (AI) in transforming human resource management in various fields, especially the recruitment and selection. Almost half of the evaluated studies, as the analysis showed, were related to recruitment and indicates the popularity of automated resume screening, natural language processing, and machine learning classifiers in talent acquisition (Kamdar et al., 2024). This supremacy is not particular, since the recruitment processes produce massive datasets that can be analyzed with the help of algorithms and the organization is pressured to maximize its hiring efficiency. Nevertheless, the fact that the studies of this field are concentrated, begs the question of the relative inattention to other HR functions, including training, development, and retention, which are also vital to the inclusivity of the workforce (Triana et al., 2021). The research agenda should be balanced to make sure that AI-driven innovation serves the entire employee life-cycle and not overemphasize the entry points.

Algorithmic bias as a structural challenge

One of the themes that appear to be consistent in the literature is the bias that still exists in the systems of AI. The use of recruitment algorithms that treat candidates who have non-standard resumes unfairly, performance appraisal frameworks that underestimate divergent communication patterns, and promotion systems that reenact historical injustices all underline the extent to which algorithmic discrimination is embedded in structural inequities. This implies that AI in HRM is non-neutral but rather the mirror of the organizational data and decision-making legacies (Rodgers et al., 2023). Noteworthy, this structural issue brings into question the belief that technological means can provide objectivity in the human resources activities by itself (Cross and Swart, 2022). Rather, the introduction

of AI demands a pivotal understanding of how systemic discrimination is coded and reproduced by algorithmic design.

Effectiveness and limitations of mitigation strategies

Another point that came out during the review was that most of the mitigation strategies used in technical bias are mostly in-processing where algorithms are directly altered to allow fairness. Although these strategies proved to be very effective in terms of predictive accuracy and fairness, they have also complicated transparency and interpretability as a result of making the model more complex. Balancing techniques like data balancing proved mediocre yet risked simplification of diversity into real-world (Zlobin and Bazylevych,). Strategies that relate to governance, such as human-in-the-loop auditing, were less prevalent yet found great approval in terms of promoting accountability (Häußermann, & Lütge, 2022). This unbalanced adoption underscores a conflict between technical convenience and ethical accountability: whereby organizations are induced to technical resolutions, long-term fairness necessitates cultural and structural adjustments within governance systems.

Quality and rigor in fairness research

The quality evaluation uncovered that over 40 percent of the studies were medium-quality in nature and in most cases were deficient in transparency in terms of dataset disclosure, and fairness measures. This lack of balance obstructs comparability and reproducibility of findings and the possibility of making powerful generalizations. The fact that fairness in HRM is not measured in a standardized way, and the studies searched the measures in various ways, using either the disparate impact ratios or equal opportunity measures or creating their own indicators (Kuliyev et al.,). The urgent necessity to enhance the evidence base is to methodologically standardize, report transparently, and focus more on open datasets in order to reproduce and cross-verify findings.

Implications for inclusive workforce practices

One of the implications of the findings is that algorithmic fairness needs to be incorporated into more extensive diversity, equity, and inclusion (DEI) programs in organizations. When designed and managed properly, AI tools are not to be seen as the tools of efficiency promotion but as the means to enforce inclusivity. Such integration highlights the importance of a human-centered approach in which AI should and will not replace the position of ethical HR decisions.

Future directions for research and practice

The current review identifies some of the opportunities in future research. To begin with, areas of under-researched topics including the issue of training, development, and retention require more academic focus to create a complete picture of the role of AI in HRM. Second, cross-sectoral and cross-cultural research is required to understand how the issue of algorithmic fairness differs according to the organizational and regulatory settings. Third, hybrid models of integrating technical mitigation and organizational governance are to be constructed and experimented as to their long-term effectiveness. Lastly, computer scientists, HR practitioners, ethicists and policymakers must work together interdisciplinarily to come up with frameworks that would balance the effectiveness of algorithms with human considerations of fairness.

Towards responsible AI in HRM

Finally, the conclusions of this review underline that algorithmic fairness in HRM is not the absolute goal but an on-going process that needs to be monitored, adjusted, and addressed. Although AI technologies present huge opportunities in streamlining human resource practices, it is associated with

potential risks to perpetuate inequalities in case it is applied in an unequal manner without ensuring fairness. To resolve these risks, there should be a paradigm shift: it is necessary to shift towards a strategy that is less reactive in bias correction, and more proactive in terms of fairness-by-design, which will incorporate inclusiveness at each phase of the AI lifecycle. Combining technical, ethical, and organizational approaches, HRM may become a sphere where AI will reinforce the efficiency and equity and workforce practices should be aligned to the principles of diversity and inclusion.

Conclusion

Based on this review, it is revealed that although Artificial Intelligence is now an extremely effective instrument in the contemporary Human Resource Management, its implementation poses fundamental questions of fairness, inclusiveness, and responsibility. The technical bias-mitigation strategies, especially the in-processing strategies have proven promising, though they are not enough on their own since it is not possible to have algorithmic fairness without governance, transparency, and inclusive organizational practices. Enhancing equity within AI-driven HRM does not merely presuppose proper methodological standards and fairness-by-design principles but also the more comprehensive cultural outlook on diversity, equity, and inclusion. When IT transformations are balanced with the human-centered vision, companies can no longer be satisfied with efficiency improvement and develop the trust, enhance the diversity of the workforce, and foster truly inclusive traditions. The only opportunity to go is, ultimately, to combine technical solutions with ethical governance, which will make AI the enforcer of fairness but will not enforce systemic workplace injustices.

References

- [1] Ayibam, J. N. (2024). Adaptive Governance for AI Companies: Resolving the OpenAI Leadership Crisis through Dynamic Stakeholder Prioritization. *Alkebulan: A Journal of West and East African Studies*, 4(2), 40-53.
- [2] Benabou, A., Touhami, F., & Demraoui, L. (2024, May). Artificial intelligence and the future of human resource management. In *2024 International Conference on Intelligent Systems and Computer Vision (ISCV)* (pp. 1-8). IEEE.
- [3] Cheong, B. C. (2024). Transparency and accountability in AI systems: safeguarding wellbeing in the age of algorithmic decision-making. *Frontiers in Human Dynamics*, 6, 1421273.
- [4] Cross, D., & Swart, J. (2022). The (ir) relevance of human resource management in independent work: Challenging assumptions. *Human Resource Management Journal*, 32(1), 232-246.
- [5] Häußermann, J. J., & Lütge, C. (2022). Community-in-the-loop: towards pluralistic value creation in AI, or—why AI needs business ethics. *AI and Ethics*, 2(2), 341-362.
- [6] Kamdar, I., Pandey, P., & Singh, A. (2024). Talent Acquisition And Resource Management Using Natural Language Processing.
- [7] Kumari, D. A., Shaikh, I. A. K., Gangadharan, S., Kiran, P. N., Ramya, S. R., & Nargunde, A. S. (2024, April). Machine Learning for Diversity and Inclusion: Addressing Biases in HR Practices. In *2024 5th International Conference on Recent Trends in Computer Science and Technology (ICRTCST)* (pp. 182-187). IEEE.
- [8] Malik, A., Budhwar, P., Mohan, H., & NR, S. (2023). Employee experience—the missing link for engaging employees: Insights from an MNE's AI-based HR ecosystem. *Human Resource Management*, 62(1), 97-115.
- [9] Rodgers, W., Murray, J. M., Stefanidis, A., Degbey, W. Y., & Tarba, S. Y. (2023). An artificial intelligence algorithmic approach to ethical decision-making in human resource management processes. *Human resource management review*, 33(1), 100925.

- [10] Triana, M. D. C., Gu, P., Chapa, O., Richard, O., & Colella, A. (2021). Sixty years of discrimination and diversity research in human resource management: A review with suggestions for future research directions. *Human Resource Management*, 60(1), 145-204.
- [11] Vivek, R. (2023). Enhancing diversity and reducing bias in recruitment through AI: a review of strategies and challenges. *Информатика. Экономика. Управление/Informatics. Economics. Management*, 2(4), 0101-0118.
- [12] Vohra, N., Chari, V., Mathur, P., Sudarshan, P., Verma, N., Mathur, N., ... & Gandhi, H. K. (2015). Inclusive workplaces: Lessons from theory and practice. *Vikalpa*, 40(3), 324-362.

Authors:

Author 1:

Full name: Sungida Akther Lima

Email: sungidaakther.lima@gmail.com

Affiliation: MBA | Washington University of Science and Technology

ORCID iD: <https://orcid.org/0009-0003-0449-3340>

Google Scholar Profile: [Sungida Akther Lima](#)

Brief biography: Sungida Akther Lima is an experienced Human Resource Manager with a passion for workforce development, Sungida specializes in recruitment, employee engagement, and organizational compliance. She has published five research papers and presented one conference paper in the field of HRM, contributing valuable insights to the discipline. Her work blends strategic HR planning with compassionate leadership, aiming to create inclusive and high-performing workplaces. Sungida is committed to continuous learning and driving impactful change in human resource practices.

Author 2:

Full name: Md Mostafizur Rahman

Email: mostafizz121296@gmail.com

Affiliation: IMIS | South Westphalia University of Applied Sciences, Germany

Brief biography: Md. Mostafizur Rahman is an accomplished HR professional and researcher with an MBA and BBA in Human Resource Management and Management Studies from Noakhali Science and Technology University. Currently serving as an HR Officer at ABUL KHAIR GROUP, he specializes in optimizing HR operations, reducing employee turnover, and enhancing organizational efficiency. With expertise in recruitment, payroll, and employee relations, he holds certifications in KPI, Labour Law, and HRM. As a dedicated researcher, he has published journals in HR, workforce analytics, and organizational behavior. His work reflects his commitment to advancing organizational efficiency and employee well-being through evidence-based strategies. His dedication to continuous learning and professional development is evident in his pursuit of certifications and ability to adapt to the evolving HR landscape. With a strong academic background, practical experience, and a passion for research, Md. Mostafizur Rahman is a promising professional and researcher in Human Resource Management, making valuable contributions to the HR community.