

Exploring Ethical Considerations:
Generative AI's Impact on Current & Future HR Practices

Eliana Brereton

PSYC*4880: Undergraduate Honours Thesis, W24

Written under the supervision of Dr. Joshua (Gus) Skorborg

Table of Contents

<i>Abstract.....</i>	<i>2</i>
<i>Introduction.....</i>	<i>3</i>
<i>Methodology.....</i>	<i>8</i>
Qualitative Synthesis & Meta-Ethnographic Review	9
Inclusion/Exclusion of Research Articles	10
Translation of Themes & Key Insights	12
<i>Literature Review.....</i>	<i>14</i>
Benefits of GAI to HRM Professionals.....	16
Streamlining of Recruitment Activities.....	16
GAI-Powered Organizational Chatbots.....	18
Large Language Models.....	19
Ethical Considerations of GAI Usage in HR Contexts	22
Contaminated Training Data & Biased Decision Making	23
Copyrighted Content & GAI.....	27
Privacy Concerns	28
Explainability, Accountability, and Trust of Algorithmic Decisions	30
<i>Conclusion</i>	<i>33</i>
<i>References.....</i>	<i>35</i>

Abstract

This research study investigates the currently evolving ethical landscape of Human Resource practices, marked specifically by the increasing influence of Artificial Intelligence (AI), and the recent surge of development in Generative AI (GAI) technologies. The shift from traditional Human Resources (HR) methods to AI-assisted processes, GAI in particular, has redefined the “norm” of many HR practices, and likely will continue to do so as the technology develops further. Organizations have historically increased their adoption of AI-driven tools to optimize their HR practices, and these tools do offer many benefits to HRM professionals.

However, the omnipresent and increased use of AI and GAI in Human Resource Management (HRM) practices also raises significant ethical concerns, ranging from algorithmic discrimination to privacy issues, and many more. This thesis aims to address critical gaps that exist in the currently published literature on this topic, by exploring the benefits as well as ethical considerations of this technological adoption in HR. The current literature lacks a comprehensive understanding of the ethical ramifications & surrounding perceptions of GAI’s widespread adoption in HR practices. Bridging these gaps and contributing to this higher discussion is vital to ensure fair, ethical, and transparent HR practices that benefit organizations and employees alike going forward.

This thesis employs a qualitative literature analysis and follows a meta-ethnographic methodology, building on existing frameworks & literature. By adopting this approach, this thesis contributes a more nuanced understanding of the ethical implications & trustworthiness perceptions of GAI in HRM contexts and helps lay the groundwork for further exploration and intervention in this research area. This thesis not only acknowledges the potential benefits of GAI in HR, but also highlights ethical considerations and pitfalls, as well as the need for collaborative efforts across disciplines to guide its responsible implementation in HR contexts.

Introduction

In the rapidly changing landscape of Human Resource (HR) practices, the recent infusion of Artificial Intelligence (AI) has completely reimagined how employees perform their work tasks. Machine Learning (ML) and earlier iterations of AI have become commonplace in conducting many internal HR processes over recent decades, with the recent prevalence of hiring algorithm adoption exemplifying this increased popularity (Yam & Skorburg, 2021). The widespread adoption of novel AI technology in HR, coupled with the specific advancements & technological capabilities they possess, signifies a discipline-wide paradigm shift away from how tasks are completed in HR in modern times. AI-assisted processes now dominate many aspects of HR practices, having become the new “norm” (Li et al., 2021).

A critical research knowledge gap exists in comprehending the benefits as well as the ethical, philosophical, and psychological implications associated with the use of Generative AI in HR practices. Despite ongoing research on AI's application in HR practices, there remains a dearth of studies examining the ethical implications (and psychological perceptions) of generative AI becoming a “norm” within HR practices. This lack of current literature is most likely fueled by AI's ever-evolving nature and exponential development in recent years.

Many organizations have embraced AI-driven tools like applicant tracking systems, chatbots, and recommender systems in an attempt to streamline their HR practices (mainly in recruitment aspects) for decades (Bogen & Rieke, 2018). In the last ten years, organizations have increased their adoption of these AI-powered technologies (like hiring algorithms), being attracted by the potential for broader candidate pools, reduced recruitment expenses, and diminished human bias (Yam & Skorburg, 2021). This rapid infusion of AI into HR, however, has also raised significant ethical concerns, particularly regarding issues such as contaminated model training data, bias and issues of ethicality, privacy, copyright, transparency concerns, and more.

Neural network algorithms and other generative models have been shown to exacerbate human biases in certain contexts, further perpetuating disparities between minority and privileged groups (Pagano, 2023). Additional research has shown that GAI tools may also

demonstrate effects like sampling biases and negative set biases (Ooi et al., 2023). However, many organizations have adopted this technology because of “reduced human bias”, as aforementioned. This obvious conflict in the literature/public opinion calls for a reassessment of how we develop and employ these systems in HR contexts, aiming to better eliminate any human rights or distributive justice violations. Given the substantial impact HR decisions can have on individuals' well-being (Bogen & Rieke, 2018), the utilization of cutting-edge Generative AI (GAI) in HR demands further ethical scrutiny, to prevent inadvertent potential harm.

The emergence of GAI represents a notable and unique advancement in AI; it introduces generative capabilities that go far beyond the capabilities of predictive models of the past (Bommasani et al., 2021). GAI presently finds applications across various domains, in many disciplines, including natural language processing and autonomous systems (Wach et al., 2023). As the models themselves and their abilities have evolved, their role in shaping how crucial HR decisions are made has followed suit. Despite GAI's massive potential to optimize and automate many HR practices, questions regarding its ethical implications persist (Dennis & Aizenberg, 2022, Jarrahi et al., 2021).

Currently existing research indicates that positive applications of GAI in HRM practices are plentiful. Recent literature asserts that GAI holds the potential to optimize certain HR practices. This includes the optimization of the recruitment process, as well as providing HRM professionals with resources such as Chatbots and LLMs to assist in and speed up their day-to-day tasks. GAI also hosts the capability to assist HRM professionals in training and development initiatives, influencing resource allocation outcomes, and eventually bettering employee engagement (Ooi et al., 2023). Harnessing the technological capabilities of GAI within HR practices could also potentially entail employees to spend less time on laborious tasks that could be automated via GAI, which would (theoretically) result in more time allocated for other more strategic or useful tasks (Raj et al., 2023).

The widespread organizational use of AI generally in HR practices, however, has also historically accompanied many ethical dilemmas and questions related to bias & algorithmic

discrimination, copyright concerns, as well as privacy and transparency issues (Tambe et al., 2019, Dennis & Aizenberg, 2022).

The ethical discourse surrounding the use of AI in HR practices extends beyond just practical considerations; it delves into fundamental philosophical questions about justice, autonomy, and human agency. The growing dependence on AI in HR also produces ethical considerations specifically related to the perceived trustworthiness of the decision-making processes of these tools (Newman & Harmon, 2020). The newly adapted use of GAI, with its significantly superior nature and faster task completion compared to previous AI iterations, is anticipated to give rise to a host of new ethical dilemmas.

The results of decisions made in Human Resource Management (HRM), which encompass aspects such as hiring, performance evaluations, project preferences, and even terminations, have profound implications for individuals, organizations, and society (Newman & Harmon, 2020). People's livelihoods often depend on their accessibility to work and generate an income for themselves. The potentially nefarious ethical implications these technologies hold give rise to concerns about fairness and discrimination when they are deployed in HR contexts (Ooi et al., 2023). As GAI increasingly infiltrates itself into commonplace HR practices, the potential impact on individuals' lives and career trajectories cannot be overstated.

This thesis aims to contribute to the current body of research investigating the ethical considerations surrounding the utilization of Generative AI within HR practices. To do so, this thesis reviews and analyzes the currently existing literature surrounding this topic, to highlight and examine both potential opportunities and potential risks of this recent technological adoption. This thesis' aim is primarily to capture the currently existing understanding of this research area to better inform HR professionals, ML developers, and relevant stakeholders alike when making use of these models in HR contexts.

This thesis is primarily a qualitative literature synthesis in nature, and mostly extracts overarching themes from a set pool of currently existing hand-picked literature. Thus, a majority of the thesis is descriptive in nature and reflects the current public opinion/standing on this topic. However, in the conclusion of this thesis, I do try to include some of my own opinion and more

normative/evaluative information to potentially help better inform relevant stakeholders of all perspectives.

The overarching and eventual goal is to contribute to the development of specific ethical frameworks and organizational guidelines for the development and deployment of GAI-powered technologies within HR contexts, to reduce individual and organizational instances of harm. Guided by principles of human rights (Yam & Skorburg, 2021) and procedural/distributive justice principles (Miller, 2017), we must attempt to demystify and better understand the potential biases and harms that could come alongside the usage of GAI in HR practices.

Furthermore, this thesis also strives to acknowledge that the psychological perceptions of both organizations and employees regarding the trustworthiness of algorithmic decisions introduces new and important layers of complexity to the ethical discourse. These aspects are largely unaddressed by the currently existing literature. The precise ethical, philosophical, and psychological ramifications of AI in HR processes remain inadequately understood, and this is accompanied by a lack of comprehensive ethical frameworks in this area (Tambe et al., 2019). The existing frameworks also often neglect the employee perception aspect of these tools' applications, as highlighted by Köchling and Wehner (2020).

Additionally, there is a noticeable absence/lack of interdisciplinary collaboration within research in this space, concerning the development and deployment of GAI technologies within HR contexts. I/O psychologists, HRM professionals, and Machine Learning developers, each playing a crucial role in influencing these potential ethical outcomes, need to collaborate to reduce potential negative outcomes with future use of AI in these sensitive contexts (Gonzalez et al., 2019). This collaboration is imperative to guide the development and deployment of these technologies towards morally just outcomes.

Bridging these gaps between industry professionals and technological/human factors is vital for ensuring fair, ethical, and transparent Generative AI-driven HR practices. A collaborative effort is essential to navigate the complexities and mitigate potential biases that exist in the intersection of technology, human perception, and ethical considerations of AI in HR

(Gonzalez et al., 2019), GAI included. The significance of addressing this identified knowledge gap lies in its critical role in steering society towards fair, ethical, and transparent HR practices with the widespread, normalized use of Generative AI technologies within these practices.

Advocating for the regulation of the use of GAI in HR becomes essential to establish a level playing field, foster fair competition, protect intellectual property rights, and safeguard privacy (Wach et al., 2023). Legal justifications, encompassing universal human rights laws, as well as the APA core principles of fidelity, responsibility, and justice, also emphasize the moral imperative to align GAI practices with ethical standards and human rights (Yam & Skorbun, 2021; Landers & Behrend, 2023; American Psychological Association, n.d.).

Methodology

In this study, I conduct a comprehensive methodological approach, drawing inspiration from the work of Hunkenschroer and Luetge's paper, "Ethics of AI-Enabled Recruiting and Selection: A Review and Research Agenda" (2022). This type of analysis involves conducting a hybrid approach of a qualitative literature synthesis and a meta-ethnographic review of the existing literature in your topic of focus. I chose to implement this methodology for a myriad of reasons, particularly though, due to its capability to address larger, open-ended research questions. Since this thesis is particularly focusing on the ethical considerations, including both the benefits and drawbacks of Generative AI (GAI)-enabled Human Resources (HR) practices, I felt this methodology was most suitable to address the research questions of this thesis.

Hunkenschroer and Luetge (2022) conducted a review of the ethicality of AI enabled recruiting in four stages, including analyzing and reviewing how it is currently discussed in existing literature, categorizing said literature according to different assumed perspectives, mapping out ethical considerations (found in the forms of opportunities, risks, and ambiguities), and describing approaches to potentially mitigate ethical risks in practice. This thesis was approached and conducted with a very similar research methodology to address the research topic(s) at hand.

Once determining my topic of focus, I sought to discover more about the current understanding of GAI in HR as discussed in the currently existing literature. I then categorized the literature into themes, and highlighted commonalities and differences, to eventually contribute further to the collective understanding of the research topic at hand. In my thesis, I map out ethical considerations in the form of both potential opportunities and ethical risks of GAI in HR, as well as mentioning potential approaches to mitigating risks. The overarching goal of employing such a methodology in this thesis was to shed light and hopefully inspire further discussion in the area of ethicalities of GAI in HR. This decision and specificity of research topic was chosen due to the currently vague understanding of this concept across disciplines, as well as the lack of collaborative frameworks/guidelines surrounding GAI in HR.

Qualitative Synthesis & Meta-Ethnographic Review

A qualitative literature synthesis can be defined as a methodology in which study findings are systematically interpreted through a process of calculated judgements and comparisons, to represent a deeper theme within the larger pool of existing literature (Sattar et al., 2021). In most qualitative literature synthesis, the findings of other qualitative studies are pooled and discussed (Sattar et al., 2021, Barnett-Page & Thomas, 2009). Certain types of qualitative synthesis allow for the inclusion of quantitative research studies alongside qualitative studies (Barnett-Page & Thomas, 2009), which is what I have done whilst forming this thesis. By reviewing and integrating these diverse sources, a qualitative synthesis can shed light on specific questions or phenomena and offer deeper insights and explanations that a single study alone may not be able to provide (Hannes & Macaitis, 2012).

The meta-ethnographic approach, as outlined by Noblit & Hare (1988, 2019), serves as a guiding framework for this methodology. Meta-ethnography is another method for the synthezation of qualitative studies. It initially involves the identification and review of studies related to a specific phenomenon of interest (Noblit & Hare, 1988). These studies are systematically examined, leading to the reduction of more relevant studies, and a more refined specification of the phenomenon (Noblit, 2019). One of the defining features of meta ethnography is that of theory generating and pinpointing avenues for future research (France et al., 2016).

The specific steps of conducting a meta-ethnographic research study outlined by Noblit & Hare (1988, 2019) include;

1. Choosing the topic focus,
2. Deciding what is relevant to the initial interest,
3. Reading the studies repeatedly,
4. Determining how studies are related,
5. Translating the studies into one another,
6. Synthesizing translations, and
7. Expressing the synthesis in a suitable form for the audience.

Meta-ethnographic studies present several benefits, such as the capability to challenge or modify the social comprehension of a phenomenon, create models and hypotheses that can be tested, and enhance the applicability of findings to wider contexts (France et al., 2016). This methodology is also particularly beneficial to highlight currently existing deficiencies in conceptual development/literature (France et al., 2016). In contrast to other forms of review and synthesis, meta-ethnographic reviews focus on larger themes and are suitable for a wide range of literature types, including qualitative, quantitative, and conceptual works (Noblit & Hare, 1988). The flexibility of this meta-ethnographic methodology as conducted by Hunkenschroer and Luetge (2022), allowing for the inclusion of qualitative and quantitative review with a narrative synthesis element, felt the most suitable to address the specific research questions this thesis aims to address.

Inclusion/Exclusion of Research Articles

To begin the literature review and meta-ethnographic approach, I determined the focus topic by first identifying a key gap in the research area and existing literature of my chosen area of interest, which was originally AI in HR processes. Once identifying the lack of literature and research gap surrounding GAI's implications for HR practices, I completed a structured keyword-based literature search on the major online database Google Scholar.

Due to the novelty and lack of existing research on my specific area of focus, my search terms included a wide range of related topics, in hope to capture a breadth of different, and potentially valuable, perspectives on the topic (Hunkenschroer & Luetge, 2022). I combined keywords in my search inclusion criteria into 4 main categories: AI, GAI, Human Resources, and Ethics. I included articles from academic peer-reviewed journals, conference proceedings, and practitioner-oriented articles (e.g., magazine articles) that study the ethicality of GAI-powered practices in HR practices/contexts. While the searches conducted for the present thesis were not systematic, this is because this thesis aims to contribute to the broader discussion occurring regarding GAI in HR. To attempt to better identify the relevant literature to inform this thesis, I searched for articles which described both positive and negative applications/evaluations GAI being applied into HR practices. Table 1 provides an overview of the research paper collection and selection criteria.

Table 1: Criteria for Literature Search and Selection (Hunkenschroer & Luetge, 2022)

Search Terms	<i>“Generative AI”, “Bias”, “Human Resources”, “Foundation Model”, “Predictive Model”, “Large Language Model”, “GPT”, “AI”, “Artificial Intelligence”, “Hiring Software”, “Perspective”, “Narrative Analysis”, “Workplace”, “Literature Review”, “Philosophy”, “Justice”, “Fairness”, “Fairness Perceptions”, “Psychology”, “Ethics”, “AI Ethics”, “Trust Facilitation”, “Trustworthiness”, “Trustworthiness Perceptions”, “Human Rights”, “Distributive Justice”, “I/O Psychology”, “Human Resource Management”, “Recruitment”, “Framework”, “Guidelines”</i>
Search Procedure	Initial keyword search, backward search, forward search.
Language	English
Time Frame	No limitation, though generally looking for articles from 2016 onwards.
Databases	Google Scholar
Inclusion	Articles from academic peer-reviewed journals, conference proceedings, and practitioner-oriented articles (e.g., magazine articles) that study the ethicality of GAI-powered practices in HR practices/contexts., accessible in full text via open access or institutional access.
Exclusion	Articles and studies without direct relation to the ethicality of AI-enabled HR processes, letters to the editor, commentaries, interviews, reviews, conference abstracts.

Many decisions regarding inclusion and exclusion of research were made following Noblit & Hare's (1988, 2019) approach and comprehensive review. These decisions were made to try and avoid bias, further define the scope of the research area, and contribute to discussion in the research area regarding the current state of the literature (Noblit, 2019).

The main inclusion criteria consisted of articles from journals or databases, primarily related to the ethicality, philosophy, and psychology of GAI-enhanced HR practices, as well as trust perceptions of AI systems, which were accessible in full text via open access or institutional access. Articles that were found through search results were analyzed thoroughly to determine relevance to overarching themes and to determine whether they met the inclusion or exclusion criteria.

After a thorough and robust literature review, I found 40 relevant studies that contributed to the elevated discussion of GAI in HR practices that were determined worthy of inclusion into this thesis. All search results were appraised for their quality of research, potential factors of bias, and potential contribution to the study, based on guidelines outlined by Hunkenschroer and Luetge (2022) and Noblit & Hare (1988, 2019).

Translation of Themes & Key Insights

The next phase of conducting this thesis, as proposed in the methodology from Noblit & Hare (1988, 2019), as well as Hunkenschroer & Luetge (2022), involved translating the interpretive themes of each study into the others. This then enabled the generation of comprehensive insights which allowed for further discussion. All relevant articles were read, and all recommendations/insights were noted. Any disagreement between authors of papers was noted, and guidance which was based on systematic reviews of the literature rather than individual reflections was prioritized (Sattar et al., 2021).

Throughout the writing and planning process of this thesis, the main themes and differing perspectives within my chosen articles were expressed and corroborated. These themes were then synthesized into the following literature review portion of this thesis, to express the current standing of the literature in an understandable way for a variety of potential audiences (and promote future discussion).

By adopting this narrative literature review analysis and meta-ethnographic approach, my aim with this thesis is to contribute to a comprehensive and insightful examination of the ethical implications of GAI adoption in HR practices. This approach provides a solid foundation for understanding the current landscape and identifies areas for further exploration and intervention research. This ensures a robust and informed analysis of the subject matter at hand as we go forward.

Literature Review

Generative Artificial Intelligence has surged in popularity across numerous disciplines in recent years. Large Language Models (LLMs), as well as many other forms of Generative AI (GAI), have garnered significant attention from the public and in media as of late. For instance, chat.openai.com, which is host to the GPT models (OpenAI’s famously conversational set of LLMs), received a staggering 1.65B total visits as of February 2024 (SimilarWeb, 2024, OpenAI, 2023). Many people are beginning to recognize the expansive potential of these AI-powered technologies, as well as the applications they can have across varying contexts. This widespread adoption is largely attributed to GAI’s newfound ability to generate **novel** outputs based on user inputs (Bommasani et al., 2021). This has revolutionized traditional approaches to many everyday tasks for people in varying fields with access to this technology, both recreationally and for commercial/organizational purposes (Hacker et al., 2023, Ooi et al., 2023).

The cornerstone of many popular GAI applications lies in foundation models (Bommasani et al., 2021). These models are trained on vast and diverse datasets, enabling them to produce innovative (and novel) outputs based on user inputs in a variety of contexts (Bommasani et al., 2021, Ooi et al., 2023). Foundation models, often employing large-scale self-supervised learning techniques, possess the flexibility to adapt to a broad spectrum of tasks that may be different than the ones they were originally programmed for, via the process of fine tuning (Bommasani et al., 2021). GPT-4, a very popular LLM that makes use of this technology, can effectively handle a diverse array of tasks through natural language prompts, despite not undergoing explicit training for many of those tasks (OpenAI, 2023).

Despite their recent prominence, these types of models and their predictive capabilities aren't new technologies by any means—they often make use of other Machine Learning (ML) concepts, such as deep neural networks, which are well-established and have existed for many years (Bommasani et al., 2021). However, it is their newly developed generative capabilities that will be the focus of this research, as they are particularly noteworthy in both their potential positive and negative applications across many fields (Bommasani et al., 2021, Ooi et al., 2023).

The scale at which these models are being trained across billions of parameters of training data, alongside their open-source access to the public, is unlike anything we have seen before and therefore worthy of further ethical consideration.

Recent advancements in GAI have introduced **multimodal outputs**, expanding beyond just text-based outputs (Bommasani et al., 2021, OpenAI, 2023). These diverse outputs include images (e.g., DALL-E, Stable Diffusion, MidJourney), audio (e.g., AudioLM), video (e.g., Imagen Video, Phenaki), and transcription of audio files (e.g., Whisper) (Ooi et al., 2023). GPT-4, one of the GPT models, also accepts prompts and can produce outputs consisting of images, text, and audio (OpenAI, 2023). Multimodal outputs are skyrocketing in popularity.

This recent multimodal output development allows for GAI to take on a much broader spectrum of tasks and assist people in many ways it could not previously. The multimodal advancements and extremely public access of Generative AI have allowed for many people of all walks of life to be able to utilize this technology in ways that we could once only imagine. GAI is rapidly being adopted and recognized for its immense potential across many disciplines, including within Human Resources (HR). This development, while potentially bountiful, leaves us with a plethora of new ethical considerations in its wake.

Human Resource Management, as defined by Armstrong (2014), refers to a deliberate and well-planned method of overseeing an organization's most esteemed resources – its employees – who, through their individual and collaborative efforts, contribute to the fulfillment of the organization's goals. The roles of HRM professionals encompass a wide array of tasks, reflecting the diverse range of specialties and areas within the field, each requiring specific skills and focus (Armstrong, 2014).

Due to this varied nature of potential “everyday” HRM tasks, the focus of this thesis is mainly on the varied applications GAI could have on the **recruitment, selection, and onboarding** aspects of HRM (as defined by Tambe et al., 2019). Recruitment and selection decisions in HRM hold the power to influence an individual's future opportunities, including their employment prospects, income levels, residential location, and overall quality of life (Yam & Skorburg, 2021). This area of focus was chosen to try and keep a narrower research focus

overall, and in order to better cover all aspects of the potential applications to be discussed for more practical and applicable results.

In this research I hope to delve into both the positive contributions as well as some of the potential ethical challenges that may be posed by Generative AI technologies being implemented into everyday HR practices. Though there are potentially harmful ethical scenarios to consider with this technological adoption, there are also many potential positive applications of Generative AI in HR that can benefit HRM professionals and other individuals alike. It is important to explore both the positive and potential negative applications of this technology, so we can better understand it, and ensure we are not inadvertently causing harm when applying said technology in sensitive contexts (such as within HRM).

Benefits of GAI to HRM Professionals

Streamlining of Recruitment Activities

The use and integration of Machine Learning (ML) into hiring practices, historically, has involved efforts to find the best qualities for selecting candidates to fill job positions, to streamline and speed up the recruitment process for HR professionals (Garg et al, 2023). Attributes used to determine the “best” candidates span from demographic factors such as age, gender, marital status, and past annual income, to personal traits such as reaction capability, comprehensive ability, and psychological quality of potential candidates (Garg et al., 2023).

AI and GAI-based hiring algorithms can facilitate the rapid exploration of potential candidate social media profiles and user-generated content through natural language processing and social media analytics, to find such attributes and any additional information about potential candidates (Ooi et al., 2023). AI and GAI powered algorithms can also conduct high-speed mass searches on social media platforms such as LinkedIn, Facebook, and Twitter to find potential job position candidates (Yam & Skorburg, 2021). Generative AI can also assist in the process of identifying potential candidates from other various online sources, including job boards, professional networking sites, and social media platforms using services like Skim.AI (Garg et

al., 2020, skimai.com, 2024). In larger organizations that can potentially receive thousands of applications for a singular job posting, the optimization of this process holds huge potential to save HRM professionals significant amounts of time.

GAI-based algorithmic hiring tools can expand the diversity of potential job candidates through targeted promotion and matchmaking platforms (Ajunwa & Schlund, 2020, Yam & Skorburg, 2021). Recent research conducted by LinkedIn (2023) found that job posts mentioning AI or GAI specifically experienced 17% greater growth in applications over the past two years, compared to posts without such mentions, exemplifying the reach and popularity this technology withholds. Hacker's (2023) research suggests that the commonplace integration of GAI into HR practices could lead to more dynamic candidate and position matching.

GAI algorithms can also be trained to sift through resumés and identify candidates whose qualifications and experiences match the job requirements, returning results faster than human employees ever could (Ooi et al., 2023). HRM professionals can do this with services such as SkillPool, which claim to ensure that only the most “qualified individuals” are considered for further evaluation (SkillPool, 2024). This is theorized to save HR professionals significant amounts of time by reducing the manual effort needed for initial screening. Some authors theorize that this process will free up time for HR professionals to focus on more “strategic”, organizationally beneficial tasks (Raj et al., 2023).

Different GAI applications also claim to rapidly engage in reverse assessment between job posting and candidate, evaluating job relevance for potential candidates and providing strategic recommendations to recruiters (Ooi et al., 2023). One example of such services, impress.ai, claims that their platform “utilizes generative AI to enable employers to harness predictive analytics, offering valuable insights into forthcoming industry trends”. Additionally, their use of GAI facilitates the identification of job seekers possessing specialized skill sets conducive to future organizational success (impress.ai, 2024). Impress.ai also claims that the use of GAI can provide HRM professionals a more diverse talent pool by encouraging applicants from underrepresented communities to apply for positions they may not have otherwise

considered (impress.ai, 2024). Impress.ai is one of many companies that offer similar services to HRM professionals based on GAI algorithms.

The future integration of GAI into HRM practices could also facilitate the development of multimodal employee/candidate assessments. Although not yet supported by existing literature, it is conceivable that GAI-powered assessment tools could incorporate a variety of inputs, including text, voice, and visual cues, to provide a more holistic understanding of employee capabilities and potential areas for improvement. Currently existing services like HireVue already employ facial recognition software during the interview process to determine “suitable” candidates, suggesting that the adoption of other multimodal assessments in the future may not be such a far-fetched thought (HireVue, 2024).

GAI-Powered Organizational Chatbots

Another potential positive example of this adoption is the concept of GAI-powered chatbots being further implemented into the way HR professionals conduct their work tasks. Chatbots can, for example, engage with candidates in real-time, answer questions, provide information about the company and the job role, and even schedule interviews (Agunis et al., 2024). This personalized interaction can enhance the candidate experience and free up HR professionals from completing repetitive tasks manually (Agunis et al., 2024). Chatbots can also provide employees with 24/7 availability, ensuring that HR-related information is accessible at any time, which is particularly beneficial in global or remote work environments (Sebastian, 2023). A GAI-powered chatbot could handle routine inquiries from employees regarding HR policies, benefits, leave requests, etc., freeing up HR’ professionals’ time.

One example of this concept that’s currently in place is IBM’s AI chatbot called “Watson Recruitment Assistant”, used to answer potential employees’ questions, and help schedule interviews (IBM, 2024). In past years, this type of algorithmic adoption in the workplace (a real robot assistant) may have sounded like science fiction, but it is our current reality. Microsoft’s Copilot suite of AI tools, for example, hosts a range of GAI tools for use within HR contexts, including a bot that can go to meetings and take notes for you (Microsoft Copilot Studio, 2024). There has been a large spike in the popularity of such chatbots and related “helpful” GAI tools

(including Microsoft Copilot) in recent years, with a 60% increase in GAI and GAI-product mentions on LinkedIn since January 2023 (LinkedIn, 2024).

GAI-powered chatbots also host the immense potential to help make the employee onboarding process much more efficient and personalized for new employees. Chatbots can deliver personalized and interactive experiences for new hires, granting them immediate access to crucial documents, policies, and procedures (Sebastian, 2023). The chatbot could guide new hires through their onboarding experience, answering their potential questions, and ensuring a smoother transition for new employees.

By incorporating a GAI-powered chatbot into their workflow, HR professionals can streamline operations, improve employee experiences, and focus on strategic initiatives that drive organizational success. Their ability to highlight important information, address potential questions, and guide new hires through their paperwork and orientation procedures, could lead to an enhanced organizational experience, enabling HR teams to focus on higher-level tasks (Patel & Joshi, 2021).

Large Language Models

There are also many specific, practical applications of GAI via the use of LLMs specifically for HRM professionals in their everyday tasks. The natural language processing capability of LLMs make them the perfect counterpart to assist with writing-based tasks, which make up many HRM activities, such as job description creation, offer letter writing, emails, and internal communications (Agunis et al. 2024, Armstrong, 2014).

For example, a LLM like GPT-3/4 could potentially help a HRM professional create their job description/posting, draft a job offer letter to a potential candidate, or even to help craft the perfect reply to an email (Sebastian, 2023). GAI-powered tools can also help optimize the type of language used in job descriptions, to attract a diverse pool of candidates while also ensuring clarity and accuracy (Ooi et al., 2023).

Research from Raj et al. in 2023 supports the claim that the use of Generative AI poses another potential benefit to organizations by providing instant support and assistance, and operating 24/7 (a feat human employees are simply not capable of) (Raj et al., 2023). With the

integration of GAI into many HR tasks, there holds the promise of the automation of the mundane, potentially allowing HR professionals to focus on more intricate and strategic duties (Raj et al., 2023, Ooi et al., 2023). This could subsequently reduce operational expenses by minimizing the need for additional personnel (though this claim is debated in recent literature) (Raj et al., 2023). Research does show, however, that HR professionals and practitioners stand to gain a competitive advantage by leveraging the potential of GAI technologies, such as the GPT models (Raj et al., 2023, Ooi et al., 2023, Agunis et al., 2024).

The potential applications GAI can have in HR practices are still evolving, alongside the technological evolution itself. While future applications and outcomes of this technology remain uncertain, it wouldn't be unreasonable to expect the continued widespread shift towards multimodal inputs/outputs from algorithmic models, moving beyond just binary ones (Bommasani et al., 2021).

Whether positively or negatively, it is evident that Generative Artificial Intelligence (GAI) is reshaping the landscape of Human Resource Management (HRM) practices in very profound and unforeseen ways. The widespread adoption and continual advancements in GAI, particularly through Large Language Models (LLMs), have given HRM professionals many opportunities to optimize their processes and improve organizational efficiency via the use of GAI. From revolutionizing recruitment processes, to enhancing employee onboarding experiences, the integration of GAI technologies holds significant promise for streamlining HRM tasks and driving positive outcomes within organizations.

However, as we embrace these potential benefits of GAI in HRM, it's imperative to acknowledge and address the ethical considerations and potential challenges that accompany its commonplace adoption. While GAI presents exciting possibilities for innovation and efficiency, it also raises concerns about privacy, bias, and fairness in decision-making (Wach et al., 2023). The ethical implications of utilizing GAI in HRM contexts must be carefully considered to ensure responsible and equitable practices are employed, due to the sensitive and usually confidential nature of HRM tasks (Armstrong, 2014). By critically examining both the positive

contributions and potential ethical dilemmas posed using GAI in HRM, we can better understand its implications, and ensure its responsible and ethical implementation.

In the subsequent section, we will explore the ethical considerations and potential risks associated with the adoption of GAI in HRM, emphasizing the need for ethical frameworks and guidelines to guide its implementation and ensure equitable practices. Ensuring equitable practices and avoiding the unethical use of different technologies can bring value to organizations by potentially avoiding regulatory, compliance, reputational, and legal harm (Dennis & Aizenberg, 2022).

Ethical Considerations of GAI Usage in HR Contexts

The extensive integration of AI generally into HR practices has historically brought forth a myriad of ethical dilemmas, spanning across **algorithmic discrimination and bias, copyright, privacy, explainability, and accountability** issues (Dennis & Aizenberg, 2022, Yam & Skorburg, 2021). Beyond practical considerations, the ethical discourse surrounding AI's usage in HR contexts delves into fundamental philosophical questions about justice, autonomy, human agency, and the types of cultural norms that we embrace in our workplaces. The recent adoption of GAI, distinguished by its generative nature and significantly superior task completion ability, is anticipated to give rise to a host of unexpected ethical dilemmas (Wach et al., 2023, Hunkenschroer & Luetge, 2022). Thus, it is important we consider the ethical outcomes that developed from previous AI iterations being deployed in HR contexts and consider how this newfound GAI adoption could adapt or worsen the current circumstances.

In this section, I will be delving into issues surrounding bias, copyright, privacy, and trust. However, I will not be exploring further into the research areas of automation and the future of work, to keep this thesis to a reasonable length. The first section, contaminated training data and biased decision making, is a concept that is discussed most widely in the currently existing literature. It is identified as a pressing ethical dilemma across multiple prominent research studies (Dennis & Aizenberg, 2022, Tambe et al., 2019, Ooi et al., 2023). This section, as a result, is more descriptive in nature and parrots the currently echoed ethical dilemmas surrounding the use of contaminated training data for HR decision-making. Following this section and in the subsequent ones, I delve into matters of copyright, privacy, explainability, and accountability issues that may accompany the inclusion of GAI into HRM processes. I found these ethical considerations were addressed much less in currently existing literature, and thus, these sections are more evaluative and include more speculative/evaluative views as a result.

Despite ongoing research being conducted on AI's applications in HR practices, there remains a dearth of studies examining the ethical implications (and user perceptions) of GAI's recent and rapid implementation. The precise ethical and psychological ramifications of AI usage in HR processes remains inadequately understood, accompanied by a lack of

comprehensive ethical frameworks in this area of literature (Tambe et al., 2019). The existing frameworks surrounding the use of AI in HRM also often neglect the trust perception aspect of these tools' applications, as highlighted by Köchling and Wehner (2020). As GAI begins to increasingly infiltrate itself into commonplace HR practices, the potential impact on individuals' lives and career trajectories cannot be overstated. We must pre-emptively address the ethical concerns of this usage to minimize potential harm, while we still do not know the entirety of the impact the technology can have.

The results of decisions made in HR, encompassing aspects such as hiring, performance evaluations, project preferences, and even terminations, have profound implications for individuals, organizations, and society (Armstrong, 2014, Yam & Skorburg, 2021). Acknowledging the inherent bias in the organizational and widespread usage of AI, where hiring decisions hold significant consequences for individuals, underscores our moral responsibility to strive for fairness in the environments making use of these technologies (Dennis & Aizenberg, 2022, Newman & Harmon, 2020).

Generative AI, particularly in HR, presents multifaceted advantages, such as the potential to streamline processes and enhance decision-making. However, ethical concerns loom large, necessitating a more nuanced exploration of the potential outcomes in order to guide the development of frameworks and guidelines going forward.

Contaminated Training Data & Biased Decision Making

The most-discussed topic in the current literature on the ethics of AI-enabled recruiting is the occurrence of **bias** within algorithmic model training data (Hunkenschroer & Luetge, 2022). One of the core processes of generative AI (and a core process within the training of many foundation models) involves data mining, a process in which algorithms differentiate individuals based on shared traits (Barocas & Selbst, 2016). This process is highly proliferated within algorithmic hiring strategies. Throughout its lifespan, research has consistently highlighted numerous ethical concerns associated with data mining and bias when associated in organizational contexts (Tambe et al., 2019, Barocas & Selbst, 2016, Bommasani et al., 2021).

Within processes that make use of models trained on biased data, such as hiring algorithms, there exists the potential to disproportionately disadvantage members of legally protected classes, placing them at a (potentially even further) systematically relative disadvantage (Barocas & Selbst, 2016). Algorithms are not trained on a contextual, grounded understanding of our world, but instead on a small subset of training data that is often littered with bias, due to historical context that the model doesn't understand (but a human HR professional may be able to).

Foundation models don't understand *why* certain groups of people would be more or less numerous in different societal standings, nor do they possess a contextual understanding of the history that led us to said position. They are trained on their set of currently available data at a large scale (Bommasani et al., 2021). The models in question "learn" from their diverse ranges of data, which may contain biases related to race, gender, religion, or possibly other sensitive attributes (Wach et al., 2023, Bommasani et al., 2021). These biases then may be regurgitated into decision making processes that make use of such models, affecting the decision's outcome. The refinement of the predictive capabilities of these systems has been directly linked to the quantity and diversity of their training data (Bommasani et al., 2021, Wach et al., 2023). Examples of specific biases exhibited by algorithmic foundation models include sampling biases and negative set biases (Ooi et al., 2023).

When applied in HRM contexts (where decisions made by HRM professionals often have significant effects on the lives of others), it is important to consider these potentially harmful aspects of model training and data mining. Neural network algorithms and other predictive, generative models have been shown to exacerbate human biases, perpetuating disparities between minority and privileged groups (Pagano, 2023). These types of algorithms detect patterns of inequality in datasets and use them to make automated decisions that then perpetuate or exacerbate these inequalities at a large scale (Yam & Skorgburg, 2021).

If, for example, an algorithm is used in the decision making of one candidate receiving a position over another candidate, organizations could potentially face unexpected outcomes based on the model's decisions, if the model was trained on biased data. The model may choose from a

biased pool of applicants to match the current standing in society, which may reflect years of biased decisions, thus perpetuating bias in a new and poorly understood way.

The step-by-step of decisions made by hiring algorithms are also often not available to HR professionals, with only the finally selected candidates being what the HR professional sees (Zerelli, 2021). This concept has been labeled throughout the literature, on a much larger scale, as the “black box” problem in newly developed AI algorithms (Zerelli, 2021). This “black box” concept refers to the idea that the *process* in which an AI algorithm comes to its final output is not always clear or understandable to the person using the tool, and sometimes even to the developer of the model themselves. In HR contexts, this concept of “black box” algorithms being used to make decisions that hold weight on people’s wellbeing gives rise to many concerns of potential bias, as well as issues of organizational transparency & accountability.

Hua et al. (2023) affirmed in their research that gender and ethnicity diversity limitations in the development process of GPT models, as mentioned previously, are reflected within many of the model’s stages; algorithmic classification, training, and estimation, leading to inherent biases and limitations within the final/public facing models. The generation of toxic content by generative models, such as the GPT models, often stems from biased and discriminatory training data. This in turn influences prompt datasets with potentially prejudiced viewpoints (Hua et al., 2023). Moreover, the training methods that are utilized by GPT models (which are predominantly based on reinforcement learning) rely on prompt datasets labeled by humans, indicating they are still very much susceptible to developer bias (Bommasani et al., 2021, Hua et al., 2023, Dwivedi et al., 2023). The fairness and accuracy of these labeled datasets significantly impact the relevance, accuracy, and ethical integrity of the models' final responses to prompt inquiries (Dwivedi et al., 2023).

This entire process of “contaminated data” being used within GAI applications and hiring decisions can happen completely unbeknownst to the potential HRM professional utilizing said tool. This implies the need for caution when making use of such models to make sensitive HR decisions. Organizations should carefully consider their usage of AI tools, including GAI, because research indicates that perceived moral violations by AI lead to blame directed at AI itself, its developers, and the organizations utilizing such AI (Wach et al., 2023).

Despite these looming ethical concerns, organizations have historically been attracted to the use of AI in HR practices on the premise of the *lessened* human biases that these algorithms can offer (Yam & Skorborg, 2021). Much of the existing literature claims there exists the potential to reduce human bias in HR decision making with the help of AI, yet the inherent bias that lies within many of the models powering these technologies is often not mentioned within the same research (Raj et al., 2023). This conflict in the literature calls for a reassessment of how we develop and employ these systems in HR contexts, aiming to eliminate human rights/distributive justice violations.

To mitigate the potential risk of bias and discrimination in GAI usage, it is important to explore methods for increasing transparency and reducing bias in GAI, specifically within HR contexts. Organizations and HR professionals making use of GAI in their operations should try to ensure that the training data that is used to train the models powering the technologies they're making use of are diverse and representative of different populations. This could involve intentionally seeking out models that make use of datasets that encompass a wide range of demographics, backgrounds, and perspectives.

Additionally, when it comes to weight-bearing decisions such as hiring, human oversight and intervention should be a high priority when considering the incorporation of GAI and AI-assisted tools (Hunkenschroer & Luetge, 2022). HRM professionals should maintain human oversight throughout any AI decision-making process, especially in critical HR decisions such as hiring, promotions, and performance evaluations. Organizations should establish clear protocols for human intervention when AI-generated outputs raise concerns or exhibit potential bias (Hunkenschroer & Luetge, 2022).

As we discuss bias within algorithmic systems, as well as human oversight of said systems, there is one last concept I feel important to add to this section. Zerelli (2021) discusses a concept of the “double standard of algorithmic accountability” that adds an interesting layer of ethical discourse to this discussion. This argument essentially makes the assertion that human beings shouldn't necessarily be considered as “the golden standard” when it comes to transparency in providing reasoning for behaviours or decisions. We don't have access to our

own human cognition in a way that would allow for “true” transparency, or at least the same type that we sometimes ask for from algorithmic systems.

Consequently, this argument posits that a double standard exists in which sometimes we ask for a higher level of transparency from these systems than we can even consistently provide ourselves. Human beings are subject to biases and post-hoc explanations that oftentimes cloud our explanation of reasoning behind our actions (Zerelli, 2021). When we try to emulate human predictions or decisions with AI, it often means that the algorithms we create are fueled by aggregating the same intuitive (and sometimes prejudiced) human decision-making we’re trying to improve on (Zerelli, 2021). Thus, the problem of bias in AI is primarily a **human problem** at its core. This concept is an interesting one to consider when discussing the ethicalities of GAI in HR, because while it’s important that we expect a certain level of transparency from these systems, it may also be worthwhile to consider the level of which we are asking for from them.

Copyrighted Content & GAI

There is little current research/understanding surrounding the unlicensed use and imitation of copyright-protected content within certain generative/predictive models (Wach et al., 2023). Issues of **unintentional intellectual property infringement** via organizations emerge alongside organizational GAI usage (Wach et al., 2023). For organizations and HRM professionals to mitigate these risks, it is crucial to explore approaches to GAI usage that enhance trust & transparency within the Generative AI landscape.

The companies that use generative models to develop & market AI-based tools for HR professionals often do not own the material used to train said models, which questions the legitimacy of their use in organizational contexts (Peres et al., 2023; Smits & Borghuis, 2022). Legal challenges have already begun to arise due to unauthorized content usage via organizations stemming from GAI models, leading to several lawsuits (Appel et al., 2023). In a hypothetical lawsuit, if the court decided that the use of generative AI based content didn’t fit into the currently existing copyright legislation, organizations could face serious legal ramifications when deploying and/or attempting to use said models to generate content for commercial or HR based purposes (Appel et al., 2023).

Therefore, it is important for organizations and HRM professionals that make use of GAI tools to take necessary precautions to protect themselves, and ensure they are acting in compliance with legal requirements regarding copyrighted content. As much as it is important to try and act ethically regarding GAI and copyright, there is still no *real* consensus in the currently existing literature on the applicability of intellectual property rights to the content and products generated by GAI tools (Peres et al., 2023; Smits & Borghuis, 2022). This lack of consensus may exist because as a society we are still grappling with questions such as; to what extent is content produced by GAI considered original, whether or not it even can be copyrighted, and who gets credit for products generated with the use of GAI (Appel et al., 2023). The uncertainty surrounding the copyright laws and legislation of GAI-produced content should urge organizations and HRM professionals making use of such content to do so with caution and specificity.

Privacy Concerns

When many people first consider the concept of privacy, or a breach of privacy, their mind may go to stereotypical examples of police, military/governmental surveillance, personal privacy, paparazzi moments, and other instances of this nature. However, recent technological advancements, especially in the field of AI, have driven the development and recognition of many new dimensions of personal and organizational privacy across literature in the past few years.

Zerelli (2021) discusses these newfound dimensions in detail in Chapter 6 of “*A Citizen’s Guide to Artificial Intelligence*”. Zerelli emphasizes throughout the chapter that less so should we be afraid of the media-driven governmental surveillance privacy concerns of the past; but we should be instead focusing on the new privacy concerns that are arising due to the never-before-seen technological prowess of recent developments in AI technology. With different algorithmic models, companies and organizations have access to an *unprecedented* amount of data about, for example, potential job applicants (Zerelli, 2021). Algorithms can be used to make inferences about your character, and the person you are, based on a collection of data you did not necessarily intend to have included in your application or desired self-representation.

Zerelli (2021) succinctly describes this concept using a metaphorical notion of “social selves”. The concept of social selves is rooted in the idea that our identity, and really who we are, is shaped by our interaction with others. We have many differing “social selves” in many different social situations (Zerelli, 2021). Privacy plays a gigantic role in shaping and maintaining the development of our social selves. Privacy gives us the space to manage the information we share with others, thus allowing us to control the perceptions and images held by those in our social circles (Zerelli, 2021). Privacy gives us the freedom to further define our social selves by allowing us to decide what is revealed and what is concealed about our identities. With recent developments in algorithmic technology, our access and *choice* in the matter of which social selves we want to present (and where) is being suppressed, and our right to privacy violated.

In today’s digital age, the use of GAI in HR practices raises significant personal privacy concerns for both applicants and HR professionals. The deployment of AI-assisted processes in hiring and employee acquisition introduces an unethical power asymmetry between the companies employing said tools and the actual applicants that are subjected to their decision making (Yam & Skorburt, 2021). The organizations deploying said technologies hold the power to gather very specific information about potential applicants, while applicants are often left in the dark when it comes to information about the company (Yam & Skorburt, 2021, Yeung, 2018). Many applicants are often not even aware they are being subjected to such algorithmic determinism in the hiring process (Yeung, 2018).

The aggressive utilization of AI algorithms allows companies to extract vast amounts of personal data from unconventional sources such as social media and online behavior, constructing detailed digital profiles without applicants' explicit consent or awareness (Yam & Skorburt, 2021). Potential job applicants have a fundamental right to privacy, including their autonomy to control the disclosure of personal information within job applications (Yam & Skorburt, 2021). This lack of transparency and consent in algorithmically based hiring infringes upon applicants' rights to personal autonomy and self-determination and can potentially lead to unintended biases and discriminatory outcomes in hiring decisions (Tambe et al., 2019).

From an organizational perspective, the integration of GAI into HR practices also poses significant privacy risks, particularly concerning the handling of sensitive organizational data. HR professionals may inadvertently compromise organizational privacy by feeding proprietary or confidential information into generative models for decision-making purposes. The reliance on AI algorithms to process sensitive organizational data, such as employee performance metrics, salary information, and internal communications, raises concerns about data security and confidentiality (Wach et al., 2023). Inaccurate or biased outputs generated by AI models can inadvertently expose sensitive organizational information, leading to reputational damage, legal liabilities, and breaches of confidentiality agreements (Wach et al., 2023). To mitigate potential privacy risks, HRM professionals need to find a balance between leveraging GAI's capabilities for efficient decision-making while safeguarding individual privacy rights, as well as organizational confidentiality.

Explainability, Accountability, and Trust of Algorithmic Decisions

As AI technologies, particularly generative AI, become more integrated into Human Resource (HR) practices such as hiring and employee acquisition, concerns regarding the explainability, accountability, and trustworthiness of algorithmic decisions have come to the forefront of discussions within HRM.

One critical concern, cited frequently in the existing literature, is the lack of explainability that is inherent in many algorithmic decisions (Dennis & Aizenberg, 2022, Jarrahi et al., 2021, Tambe et al., 2019). While algorithmic management tools have been increasingly used across various HR functions like recruiting, scheduling, and performance monitoring, the opacity of these algorithms can lead to decisions that are difficult to understand or justify (Dennis & Aizenberg, 2022). This challenge is exacerbated because many GAI and AI algorithms are considered "black boxes", as aforementioned.

“Black box” algorithms refer to the concept that despite seeing the algorithm’s result, trying to find out more about the algorithm’s process or procedure can be a difficult (or often impossible) process (Jarrahi et al., 2021). Whether to protect intellectual property rights, or due

to a lack of technical understanding from users interacting with the model, this “black box” type of algorithmic decision-making limits transparency and hinders accountability (Jarrahi et al., 2021). When making decisions in sensitive contexts such as Human Resource Management, transparency, accountability, and explainability are all values that should be considered in every decision made (Dennis & Aizenberg, 2022, Tambe et al., 2019).

This newfound reliance on algorithmic solutions within HR also raises many new questions about accountability. Who is to be held accountable when there is an undesired hiring outcome, for example, from an algorithmic decision? The algorithm that made the decision, or the HR professional that allowed it? Traditional notions of autonomy, which are a fundamental value in employment relationships and meaningful work, have become challenged as algorithms enable more and more pervasive surveillance and extensive control over workers within organizations (Unruh et al., 2022). This shift towards algorithmic control introduces new ethical dilemmas, as algorithms can be more personalized and opaque compared to previous decision-making methods we have used in the past (Unruh et al., 2022).

Studies have also identified attitudes such as algorithm aversion among stakeholders (Jarrahi et al., 2021, Dietvorst et al., 2017). Algorithm aversion is a phenomenon that reflects a hesitation to trust algorithm-generated advice or decisions, even when algorithms may outperform humans in accuracy measures (Dietvorst et al., 2017). This lack of trust can stem from past experiences of imperfect algorithmic performance or concerns about bias and discrimination (Dietvorst et al., 2017). Moreover, the procedural character of specific algorithms may lead to their decisions being seen as authoritative by default, particularly in high-pressure work environments where workload and time pressures can result in overreliance on automated systems (Jarrahi et al., 2021). This could lead to questionable ethical outcomes based on decisions made based on AI-generated advice.

Addressing these larger challenges in AI requires a multifaceted approach (Dennis & Aizenberg, 2022). Implementing more “explainable” algorithmic solutions and creating more transparent regulations, as well as auditing mechanisms, can enhance trust and accountability in algorithmic decision-making processes (Dennis & Aizenberg, 2022). Additionally, ongoing

research is needed to understand the ethical implications of people's willingness to trust automated systems and their decisions, algorithmic management in HR practices, and the labor conditions this creates. The research from this investigation should then influence both HR professionals and GAI developers.

In this section I have covered issues spanning contaminated training data, copyright concerns, privacy issues, and explainability/accountability concerns stemming from the increased usage of GAI technologies within HR processes. While this is nowhere near an exhaustive list of potential ethical concerns, the aim of this literature review is to better refine the image of where the literature is currently standing now, to eventually aid HR professionals in making more ethically sound and research-based decisions going forward.

Conclusion

In conclusion, this exploration of the benefits and ethical implications surrounding the integration of Generative Artificial Intelligence (GAI) in Human Resource (HR) practices ultimately underscored the need for a more nuanced understanding of the subject, from relevant stakeholder across varying disciplines. Despite the rapid adoption of GAI by companies seeking its numerous advantages, due to the apparent ethical quandaries, it is imperative that we first develop more robust methods for critically evaluating the ethical ramifications before there is further widespread implementation.

Through a thorough review of the existing literature and perspectives, it becomes clear that there is no one universal ethical framework that is going to be applicable to all GAI systems, all HRM professionals, or all organizations employing such technology. The normative evaluation of usage and adoption of GAI in HR should be an ongoing endeavor, driven by interdisciplinary collaboration among I/O psychologists, HR professionals, AI developers, and organizations employing such technologies.

The phenomenon of bias exacerbated by GAI & AI in HR, while wholly concerning, should not necessarily lead us to a blanket rejection of its use overall. Biases and cognitive heuristics are very much prevalent in human beings as well. This “double-standard” concept prompts us to perhaps question whether these algorithmic systems are inherently ***more biased*** than human decision-makers, versus whether they exacerbate biases at all. Additionally, more research needs to be conducted to create more concrete guidelines as to when/how these technologies should be applied most effectively, especially in sensitive HR contexts.

It is crucial to acknowledge, however, that perfection is unattainable, both in human decision-making and in GAI/AI systems. Instead of perfection, the focus should be on whether GAI represents an **improvement** over current practices, or if its increasing popularity will lead to unforeseen ethical dilemmas. This necessitates a more **case-by-case** evaluation of GAI’s implementation, and continuous monitoring from HRM professionals to mitigate negative ethical outcomes to their greatest extent possible. The beginning of this process involves HRM

professionals becoming more **aware** of the potential risks that are posed by incorporating GAI into their work. While the potential benefits of GAI in HR are undeniable, we must proceed with caution due to identified risks and the ongoing emergence of potential future issues.

Ultimately, the most responsible and ethical integration of Generative AI in HR practices will require a balanced approach that prioritizes **fairness, transparency, and ongoing evaluation**. By better equipping ourselves in going forth through the complex interplay of technological advancements and ethical considerations, we can strive towards a future where GAI contributes positively to organizational decision-making, while still upholding ethical standards and fostering trust.

References

- Aguinis, H., Beltran, J. R., & Cope, A. (2024). How to use Generative AI as a human resource management assistant. *Organizational Dynamics*, 101029. <https://doi.org/10.1016/j.orgdyn.2024.101029>
- AI features for teams and Classic Bots (contains video) - microsoft copilot studio. Microsoft Copilot Studio | Microsoft Learn. (n.d.). <https://learn.microsoft.com/en-us/microsoft-copilot-studio/advanced-ai-features>
- AI resume screening: Find your perfect hire. SkillPool. (2024, February 6). <https://skillpool.tech/>
- Ajunwa, I., & Schlund, R. (2020). Algorithms and the social organization of work. *The Oxford Handbook of Ethics of AI*. <https://doi.org/10.1093/oxfordhb/9780190067397.013.52>
- American Psychological Association. (n.d.). Ethical principles of psychologists and code of conduct. American Psychological Association. <https://www.apa.org/ethics/code>
- Amnesty International. (2019). Surveillance Giants: How The Business Model of Google and Facebook Threatens Human Rights. Amnesty International.
- Armstrong, M. (2014). *A Handbook of Human Resource Management Practice, 13th Edition*. Kogan Page.
- Bankins, S. (2021). The ethical use of Artificial Intelligence in human resource management: A decision-making framework. *Ethics and Information Technology*, 23(4), 841–854. <https://doi.org/10.1007/s10676-021-09619-6>
- Barnett-Page, E., & Thomas, J. (2009). Methods for the synthesis of qualitative research: A critical review. *BMC Medical Research Methodology*, 9(1). <https://doi.org/10.1186/1471-2288-9-59>
- Barocas, S., & Selbst, A. D. (2016). Big Data's Disparate Impact. *California Law Review*, 104(3), 671–732. <http://www.jstor.org/stable/24758720>
- Brougham, D., & Haar, J. (2017). Smart Technology, Artificial Intelligence, robotics, and Algorithms (stara): Employees' perceptions of our future workplace. *Journal of Management & Organization*, 24(2), 239–257. <https://doi.org/10.1017/jmo.2016.55>
- Bogen, M., & Rieke, A. (2018). Help Wanted—An Exploration of Hiring Algorithms, Equity and Bias. Upturn. <https://www.upturn.org/static/reports/2018/hiring-algorithms/files/Upturn%20--%20Help%20Wanted%20-%20An%20Exploration%20of%20Hiring%20Algorithms,%20Equity%20and%20Bias.pdf>

- Bommasani, R., Hudson, D. A., Adeli, E., Altman, R., Arora, S., von Arx, S., Bernstein, M. S., Bohg, J., Bosselut, A., Brunskill, E., Brynjolfsson, E., Buch, S., Card, D., Castellon, R., Chatterji, N. S., Chen, A. S., Creel, K. A., Davis, J., Demszky, D., et al. (2021). On the Opportunities and Risks of Foundation Models. ArXiv. Retrieved from <https://crfm.stanford.edu/assets/report.pdf>
- Bonnefon, J.-F., Rahwan, I., & Shariff, A. (2023). The Moral Psychology of Artificial Intelligence.
- Chat.openai.com Website Traffic Monitoring. (2024).
<https://www.similarweb.com/es/website/chat.openai.com/>
- Chatbot for HR - IBM watsonx assistant*. IBM. (n.d.). <https://www.ibm.com/products/watsonx-assistant/human-resources>
- da Motta Veiga, S. P., Figueroa-Armijos, M., & Clark, B. B. (2023). Seeming ethical makes you attractive: Unraveling how ethical perceptions of ai in hiring impacts organizational innovativeness and attractiveness. *Journal of Business Ethics*, 186(1), 199–216.
<https://doi.org/10.1007/s10551-023-05380-6>
- Dennis, M. J., & Aizenberg, E. (2022). The ethics of AI in human resources. *Ethics and Information Technology*, 24(3). <https://doi.org/10.1007/s10676-022-09653-y>
- Dietvorst, B. J., Simmons, J. P., & Massey, C. (2018). Overcoming algorithm aversion: People will use imperfect algorithms if they can (even slightly) modify them. *Management Science*, 64(3), 1155–1170. <https://doi.org/10.1287/mnsc.2016.2643>
- Dwivedi, Y. K., et al. (2023). "So what if ChatGPT wrote it?" Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy. *International Journal of Information Management*.
- France, E. F., Wells, M., Lang, H., & Williams, B. (2016). Why, when and how to update a meta-ethnography qualitative synthesis. *Systematic Reviews*, 5(1). <https://doi.org/10.1186/s13643-016-0218-4>
- Garg S, Sinha S, Kar AK, Mani M. A review of machine learning applications in human resource management. *Int J Product Perform Manag*. 2022; 71(5):1590–610. doi:10.1108/IJPPM-08-2020-0427.
- Gonzalez, M., Capman, J., Oswald, F., Theys, E., & Tomczak, D. (2019). “Where’s the i-o?” Artificial Intelligence and machine learning in talent management systems. *Personnel Assessment and Decisions*, 5(3). <https://doi.org/10.25035/pad.2019.03.005>
- Guidance on the conduct of narrative synthesis in systematic reviews. (n.d.).
<https://www.lancaster.ac.uk/media/lancaster-university/content-assets/documents/fhm/dhr/chir/NSsynthesisguidanceVersion1-April2006.pdf>

- Hacker, P., Engel, A., & Mauer, M. (2023). Regulating CHATGPT and other large generative AI models. *2023 ACM Conference on Fairness, Accountability, and Transparency*. <https://doi.org/10.1145/3593013.3594067>
- Hannes, K., & Macaitis, K. (2012). A move to more systematic and transparent approaches in qualitative evidence synthesis: Update on a review of published papers. *Qualitative Research*, 12(4), 402–442. <https://doi.org/10.1177/1468794111432992>
- Home. impress.ai. (2024, February 24). <https://impress.ai/>
- Hua, S., Jin, S., & Jiang, S. (2023). The limitations and ethical considerations of chatgpt. *Data Intelligence*, 1–38. https://doi.org/10.1162/dint_a_00243
- Hunkenschroer, A. L., & Luetge, C. (2022). Ethics of ai-enabled recruiting and selection: A review and research agenda. *Journal of Business Ethics*, 178(4), 977–1007. <https://doi.org/10.1007/s10551-022-05049-6>
- Jarrahi, M. H., Newlands, G., Lee, M. K., Wolf, C. T., Kinder, E., & Sutherland, W. (2021). Algorithmic management in a work context. *Big Data & Society*, 8(2), 205395172110203. <https://doi.org/10.1177/20539517211020332>
- Köchling, A., & Wehner, M. C. (2020). Discriminated by an Algorithm: A systematic review of discrimination and fairness by algorithmic decision-making in the context of HR recruitment and HR development. *Business Research*, 13(3), 795–848. <https://doi.org/10.1007/s40685-020-00134-w>
- Landers, R. N., & Behrend, T. S. (2023a). Auditing the AI auditors: A framework for evaluating fairness and bias in high stakes AI predictive models. *American Psychologist*, 78(1), 36–49. <https://doi.org/10.1037/amp0000972>
- Li, L., Lassiter, T., Oh, J., & Lee, M. K. (2021). Algorithmic hiring in practice. *Proceedings of the 2021 AAAI/ACM Conference on AI, Ethics, and Society*. <https://doi.org/10.1145/3461702.3462531>
- LinkedIn. *Future of work report*. AI at Work. (2023, November). <https://economicgraph.linkedin.com/research/future-of-work-report-ai>
- Machine Learning and artificial intelligence solutions*. Skim AI. (2024, February 24). <https://skimai.com/>
- Miller, D. (2017). Distributive justice. *Distributive Justice*, 297–320. <https://doi.org/10.4324/9781315257563-10>
- Newman, D. T., Fast, N. J., & Harmon, D. J. (2020). When eliminating bias isn't fair: Algorithmic reductionism and Procedural Justice in human resource decisions. *Organizational Behavior and Human Decision Processes*, 160, 149–167. <https://doi.org/10.1016/j.obhdp.2020.03.008>

- Noblit, G. W. (2019). Meta-ethnography in Education. Oxford Research Encyclopedia of Education. <https://doi.org/10.1093/acrefore/9780190264093.013.348>
- Noblit, G. W., & Hare, R. D. (1988). Meta-ethnography: Synthesizing qualitative studies. Thousand Oaks, CA: Sage Publications Ltd.
- Ooi, K.-B., Tan, G. W.-H., Al-Emran, M., Al-Sharafi, M. A., Capatina, A., Chakraborty, A., Dwivedi, Y. K., Huang, T.-L., Kar, A. K., Lee, V.-H., Loh, X.-M., Micu, A., Mikalef, P., Mogaji, E.,
- Pandey, N., Raman, R., Rana, N. P., Sarker, P., Sharma, A., ... Wong, L.-W. (2023). The potential of generative artificial intelligence across disciplines: Perspectives and Future Directions. *Journal of Computer Information Systems*, 1–32. <https://doi.org/10.1080/08874417.2023.2261010>
- OpenAI. (2023). GPT-4 Technical Report. [arXiv preprint] eprint: 2303.08774.
- Oswald, F. L., Behrend, T. S., Putka, D. J., & Sinar, E. (2020). Big Data in industrial-organizational psychology and Human Resource Management: Forward Progress for Organizational Research and Practice. *Annual Review of Organizational Psychology and Organizational Behavior*, 7(1), 505–533. <https://doi.org/10.1146/annurev-orgpsych-032117-104553>
- Pagano, T. P., Loureiro, R. B., Lisboa, F. V., Peixoto, R. M., Guimarães, G. A., Cruz, G. O., Araujo, M. M., Santos, L. L., Cruz, M. A., Oliveira, E. L., Winkler, I., & Nascimento, E. G. (2023). Bias and unfairness in Machine Learning Models: A systematic review on datasets, tools, fairness metrics, and identification and mitigation methods. *Big Data and Cognitive Computing*, 7(1), 15. <https://doi.org/10.3390/bdcc7010015>
- Patel, D., & Joshi, A. (2021). Artificial Intelligence and HR Chatbots in Employee Onboarding Process: A Conceptual Framework. In 2021 IEEE International Conference on Artificial Intelligence for Sustainable Future (ICAISF) (pp. 1-6). IEEE.
- Peres, R., Schreier, M., Schweidel, D., & Sorescu, A. (2023). On ChatGPT and beyond: How generative artificial intelligence may affect research, teaching, and practice. *International Journal of Research in Marketing*. Advance online publication. <https://doi.org/10.1016/j.ijresmar.2023.03.001>
- Raj, R., Singh, A., Kumar, V., & Verma, P. (2023). Analyzing the potential benefits and use cases of chatgpt as a tool for improving the efficiency and effectiveness of business operations. *BenchCouncil Transactions on Benchmarks, Standards and Evaluations*, 3(3), 100140. <https://doi.org/10.1016/j.tbench.2023.100140>
- Robert, L. P., Pierce, C., Marquis, L., Kim, S., & Alahmad, R. (2020). Designing fair AI for managing employees in Organizations: A Review, Critique, and Design Agenda. *Human–Computer Interaction*, 35(5–6), 545–575. <https://doi.org/10.1080/07370024.2020.1735391>

- Sebastian, G. (2023). Hello! This is your new HR assistant, ChatGPT! The Impact of AI Chatbots on Human Resource practices: A transformative analysis. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.4461474>
- Smits, J., & Borghuis, T. (2022). Generative AI and Intellectual Property Rights. In B. Custers & E. Fosch-Villaronga (Eds.), *Law and Artificial Intelligence: Regulating AI and Applying AI in Legal Practice* (pp. 323-344). T.M.C. Asser Press. https://doi.org/10.1007/978-94-6265-523-2_17
- Tambe, P., Cappelli, P., & Yakubovich, V. (2019). Artificial Intelligence in human resources management: Challenges and a path forward. *California Management Review*, 61(4), 15–42. <https://doi.org/10.1177/0008125619867910>
- Unruh, C. F., Haid, C., Johannes, F., & Bütte, T. (2022). Human autonomy in algorithmic management. *Proceedings of the 2022 AAAI/ACM Conference on AI, Ethics, and Society*. <https://doi.org/10.1145/3514094.3534168>
- Wach, K., Duong, C.D., Ejdy, J., Kazlauskaitė, R., Korzynski, P., Mazurek, G., Paliszkievich, J., & Ziemia, E. (2023). The dark side of generative artificial intelligence: A critical analysis of controversies and risks of ChatGPT. *Entrepreneurial Business and Economics Review*, 11(2), 7-30. <https://doi.org/10.15678/EBER.2023.110201>
- Why Hirevue? solutions for hiring at-scale online*. hirevue.com. (2024). <https://www.hirevue.com/why-hirevue>
- Yam, J., & Skorb, J. A. (2021). From Human Resources to Human Rights: Impact Assessments for hiring algorithms. *Ethics and Information Technology*, 23(4), 611–623. <https://doi.org/10.1007/s10676-021-09599-7>
- Yeung, K. (2018). A Study of the Implications of Advanced Digital Technologies (Including AI Systems) for the Concept of Responsibility Within a Human Rights Framework (SSRN Scholarly Paper ID 3286027). Social Science Research Network. <https://papers.ssrn.com/abstract=3286027>.
- Zerilli, J., Danaher, J., Maclaurin, J., Gavaghan, C., Knott, A., Liddicoat, J., & Noorman, M. E. (2021). In *A citizen's Guide to Artificial Intelligence*. The MIT Press.