

APPLICANT REACTIONS TO AI AUTOMATED VIDEO INTERVIEWS: THE ROLE
OF INTERVIEW SELF-EFFICACY IN PROCEDURAL JUSTICE PERCEPTIONS

JACOB A WATSON

Bachelor of Science in Psychology

Ball State University

August 2015

submitted in partial fulfillment of requirements for the degree

MASTER OF ARTS IN PSYCHOLOGY

at

CLEVELAND STATE UNIVERSITY

MAY 2024

We hereby approve this thesis for
JACOB A. WATSON
candidate for the Master of Arts in Psychology degree
for the Department of Psychology and
CLEVELAND STATE UNIVERSITY'S
College of Graduate Studies by

Committee Chair, Dr. Michael Horvath
Department of Psychology, May 15, 2024

Methodologist, Dr. Matthew Nordlund
Department of Psychology, May 15, 2024

Committee Member, Dr. Vickie Gallagher
Department of Management, May 20, 2024

Student's Date of Defense: May 15, 2024

ACKNOWLEDGEMENTS

Special Thanks

To Kelly Blair

for deploying the orbital organizing apparatus (OOA).

To Dr. Michael Horvath

for his unceasing eye for detail, superlative precision, and boundless patience.

To Dr. Matthew Nordlund

For his role as methodologist.

To Dr. Vickie Gallagher

for her role on the committee.

To Dr. Chieh-Chen Bowen

for keeping ideas practical.

APPLICANT REACTIONS TO AI AUTOMATED VIDEO INTERVIEWS: THE ROLE
OF INTERVIEW SELF-EFFICACY IN PROCEDURAL JUSTICE PERCEPTIONS

JACOB A. WATSON

ABSTRACT

This study examined how applicants perceive AI automated interviews in a vignette description and how applicants' interview self-efficacy may affect their preferences for AI or traditional interviews. The perceived procedural justice ratings and interview self-efficacy measurements were collected from undergraduate students enrolled in psychology courses at a university. Findings showed that applicants tended to rate interviews as fairer when the vignette described a job offer and when participants had higher interview self-efficacy, supporting prior research in self-serving bias. Participants generally found AI interviews to be less fair compared to traditional interviews. Hypotheses regarding interview self-efficacy on preferences for AI or traditional interviews were not supported.

Keyword: applicant reactions, AI interviews, interview self-efficacy, self-serving bias

TABLE OF CONTENTS

	Page
ABSTRACT.....	iv
CHAPTER	
I. INTRODUCTION.....	1
1.1 Artificial Intelligence.....	3
1.2 Applicant Reactions and Organizational Justice Theory.....	5
1.3 General Attitudes About AI.....	10
1.4 Applicant Reactions to Non-Traditional Selection Procedures.....	11
1.5 Self-Serving Bias.....	14
1.6 Interview Self-Efficacy.....	15
II. METHOD.....	24
2.1 Participants.....	24
2.2 Procedures.....	26
2.3 Measures.....	26
2.4 Inclusion Criteria.....	28
III. RESULTS.....	30
3.1 Testing Suitability for Control Variables.....	30
3.2 Hypothesis 1.....	36
3.3 Hypothesis 2.....	37
3.4 Hypothesis 3.....	39
3.5 Hypothesis 4.....	42
IV. DISCUSSION.....	46

4.1 Self-Serving Bias	47
4.2 The Effects of MJISE on Procedural Justice	49
4.3 Limitation	49
4.4 Future Research	51
4.5 Practical Implications	52
REFERENCES	55
APPENDICES	
A. TABLES.....	65
B. FIGURES	92
C. VIGNETTE	93
D. SURVEY.....	97

CHAPTER I

INTRODUCTION

Even before the release of ChatGPT in 2022, artificial intelligence (AI) had begun dominating the public discourse among organizations, and now the topic has become even more pronounced. The capabilities of AI tools have improved substantially in the last few years, and many speculate they will continue to improve in the near future (Thomas & Powers, 2023; Shreve et al., 2022); the adoption of AI tools in organizations has grown rapidly as the sophistication of the technology has made its use more feasible. AI's expeditious capabilities for assessing large numbers of applicants has, more relevantly, led to AI's use in selection for everything from screening to interviewing (Dilmeganni, 2022a). Tippins et al. (2021) stated that AI can now produce data on applicants' facial features and tone of voice, and controversially, that AI products claim to make valid predictive conclusions from this information. Tippins et al. expressed concern about the evidence for these claims, and there is still a great need for research on the topic of AI selection within the sphere of Industrial-Organizational (I-O) psychology more generally.

Some I-O psychologists have expressed concerns about AI use in selection and called for more research on its validity as a selection tool (Gonzales et al., 2019). Another

aspect of AI that needs further study is its impact on applicant reactions. The selection process is the first major impression many organizations will make on their applicants (Charmarro & Ahtktar, 2019). It is well established that applicants react differently to different selection procedures based on those procedures' characteristics, such as the degree of standardization or whether a procedure is digital (Folger et al., 2022).

Furthermore, applicants' reaction to a selection procedure can affect their future perceptions of that organization (Kohn & Dipboye, 1998; Latham & Finnegan, 1993). However, little research has been conducted on how applicants react to AI interviews. Initial research showing that applicants react negatively to AI selection procedures has troubling implications for organizations' applicant attraction. This makes addressing concerns over AI selection a pressing issue, as it could adversely affect the talent pool.

First, this study will replicate prior research on applicant reactions to AI selection regarding organizational justice (Hess, 2022). However, prior research has not investigated applicant characteristic moderators that could help explain individual differences in perceptions and subsequent applicant reactions to AI selection.

Organizations need to be aware of the types of people who respond more poorly to AI selection as it may have ramifications on the types of applicants whom they attract. One such characteristic may be an applicant's interview self-efficacy. There is reason to hypothesize that people with higher interview self-efficacy will rate AI interviews as less fair than traditional interviews due to self-serving bias and inadequate chance to perform, also called opportunity to perform. Therefore, this study will also test whether interview self-efficacy moderates the relationship between interview method and perceptions of procedural justice.

This literature review will discuss several topics relevant to applicant reactions to AI. First, I will introduce the general concept of AI, giving a brief background and definition. Then I will review the applicant reactions literature including related topics such as procedural justice. Afterwards, I will outline some research on general attitudes about AI before focusing on applicant reactions to AI. Finally, I will discuss how interview self-efficacy could potentially moderate applicant reactions to AI interviews and detail the corresponding hypotheses.

Artificial Intelligence

Industry experts such as Brooks (2018) have stated that AI is a vague term with various definitions, as it was adopted quickly and informally. Because of these purported inconsistencies, it is worth more precisely defining what AI for the purposes of this research, as well as an array of related items. AI comes in different forms, but the term generally refers to any computer program that attempts to act in a human way or perform a human function (McCarthy, 2007). In other words, it attempts to act intelligently. However, an AI does not necessarily need to use human means to achieve its goals. For instance, rather than using the traditional process of trial and error, an AI might instead accomplish its goals by autonomously following a pre-prescribed set of instructions written by a programmer. Some of the AI technologies available to organizations have been called cognitive technologies (Walch, 2019). These products differ from other traditional selection methods by building on predictive analyses that I-O psychologists are familiar with. AI analyses work in uncontrolled and inductive ways. That is, these products do not always follow theory driven hypothesis testing approaches and it can be difficult to tell in retrospect how an AI product has arrived at its conclusion.

Researchers' discussion of AI tools has been increasing over time (Gonzales, 2019; Hess, 2022; Asfahani, 2022). Machine learning is a method that AIs sometimes employ. The term is used to refer to an AI as being appreciably capable of learning through experience (Mitchell, 1997) and without a programmer defining many explicit rules beforehand (IBM, 2021). This recent development has allowed AI to quickly form methods of analysis to make predictions about things that organizations care about, such as job performance. Another newly popular term in the field is deep learning. Deep learning is a specific form of machine learning that has led to AI being a more widely feasible tool for analyzing data that were previously difficult for a computer to quantify, such as facial and voice recognition. These improvements have allowed AI technology to move into the business world as never before (IBM, 2021). However, it can be difficult to ascertain the method by which an AI (using machine learning or deep learning) comes to its conclusions because they are often not theory driven. This has led critics of such algorithmic solutions to refer to them as black-box solutions, though some researchers have argued that the label "heuristic" is better fitting (Cheng & Hackett, 2021). Additionally, there are currently reservations as to the predictive validity of facial and voice data within the context of selection (Tippins et al., 2021).

I-O psychologists have called for more research into AI selection procedures (Gonzales et al., 2019) as AI is becoming increasingly used in selection. To my knowledge there is no empirical investigation underway regarding the prevalence of AI selection or AI interviews, though it is the topic of many articles in the popular press, and some sources claim its use is widespread (Hsu, 2023; Gonzales et al., 2019; Stahl, 2022). Enterprise platforms claiming to use machine learning technology to select candidates are

presently available and in use (HireVue Website, 2022; Interviewer.AI website, 2022). HireVue is one such platform that offers AI tools for video interviewing, game-based assessment, and chatbots. HireVue also used facial and voice recognition software to process interviews but has since discontinued this service (Maurer, 2021). HireVue claims its products are in use by Kraft Heinz and other corporations. Overall, this suggests that AI tools are being widely used by high-profile, recognizable organizations.

Regardless of their validity at present, AI tools have the potential to substantially impact the business world and their continued adoption may now be inevitable whether I-O psychologists like it or not. But I-O psychologists are not the only party concerned with AI selection. There is evidence that some applicants may respond negatively to AI selection (Hess, 2022) and digital selection methods more generally (Folger et al., 2022). Employers utilizing AI for selection may soon find that applicants' apprehension and wariness of AI needs to be addressed. Introducing AI selection without addressing applicants' concerns may be harmful to organizations' reputations and alienate prospective applicants.

Applicant Reactions and Organizational Justice Theory

The application process affects how applicants view an organization. As will be discussed, organizations should be concerned if applicants do not feel that the selection procedures were fair. There is evidence that the perceived fairness of an organization's selection procedures affects applicants' perceptions of that organization and their intention to accept job offers from that organization. For instance, applicants who believe that an organization's selection process does not give them adequate opportunity to perform are less likely to report an intention of accepting a job offer from that

organization (Chapman et al., 2005). Additionally, the more anxious a candidate is about an AI application procedure, the less likely they are to complete the application (van Esch et al., 2021).

Applicant reactions is a topic of research that seeks to understand why applicants react the way they do to application procedures. Organizational justice is the domain of how fairness and justice are seen within any organization and is a core component of applicant reactions research. Cropanzano et al. (2007) described organizational justice as comprising three parts: distributive justice, procedural justice, and interactional justice. Distributive and interactional justice will be briefly described before focusing on Gilliland's (1993) rules for procedural justice.

Distributive justice is concerned with outcomes; people rate distributive justice as high when they feel they are receiving their "just share." Interactional justice is defined by how respectful and transparent people are in their interactions. A person may be considered interactionally just if they share appropriate information and treat others with respect and dignity.

If distributive justice is the perceived fairness of outcomes, procedural justice is the perceived fairness of how outcomes are distributed. Prior to the 1970s when Thibault and Walker (1975) introduced procedural justice, organizational justice researchers were mostly concerned with distributive justice rather than the fairness issues concerned with procedures (Lind & Tyler, 1988). Studies have since sought to explain how different selection procedures are perceived through the lens of procedural justice, and it has formed the basis of much of the applicant reactions research (Bauer et al., 2001). Applicants view some selection characteristics as less fair on average. For instance,

applicants tend to rate unstandardized procedures based on human intuition as being fairer than standardized selection methods (Highhouse, 2008).

One of the most influential frameworks in the study of applicant reactions is Gilliland's (1993) organizational justice theory of applicant perceptions which includes ten rules for procedural justice (McCarthy et al., 2017). Procedural justice can be separated into three groupings which can be subdivided into a total of ten rules (Gilliland, 1993). These groupings are formal characteristics, explanation, and interpersonal treatment.

Gilliland's (1993) formal characteristics of procedural justice include four rules: opportunity to perform, job-relatedness, consistency, and reconsideration opportunity. Opportunity to perform denotes how applicants wish to have the chance to show what they are good at. Applicants rate fairness as higher when they can steer conversations or procedures towards focusing on the characteristics in which they excel. Job-relatedness refers to how applicants prefer selection procedures that focus on job-related topics and measure features that are relevant to the job. The consistency rule denotes how applicants prefer procedures that operate the same way between applicants. Reconsideration opportunity refers to a selection system's capability to address applicant concerns or provide opportunity for appeal or redress (Leventhal et al., 1980). Reconsideration opportunity has generally been found to be more relevant for incumbent employees, such as those seeking promotion, than new applicants (Cropanzano & Wright, 2003).

Of the formal justice rules, the use of AI systems would appear more likely to violate some rules more than others. AI selection may violate the job-relatedness rule because it is based on algorithms that may discriminate on characteristics that are either

not perceived by applicants as valid, or it may discriminate based on protected class characteristics (Ferrer et al., 2020). This can happen for a variety of technical reasons, such as biased training samples and an inherently inductive style of analysis that many AI tools utilize. However, AI selection may fare well in consistency. There is some evidence that applicants may perceive AI selection as more attributionally stable than traditional selection methods (Hess, 2022) which may imply that applicants will also find AI selection to be more consistent than traditional interviews. AI procedures are not as adaptable or capable of dynamic interaction with humans. An AI procedure's ability to achieve high levels of reconsideration opportunity when used on its own may be limited as most AI are not well equipped to appraise and modify their own procedures or provide redress at this time. However, AI selection may be used as part of a greater selection system that provides other opportunities for reconsideration.

Gilliland's (1993) explanation grouping includes feedback, selection information, and honesty. Applicants prefer to get high quality constructive feedback about their performance in a prompt manner. The selection information rule refers to applicants' tendency to see selection procedures as fairer when they are given justification for why a procedure was used or why decisions were made. The honesty rule simply means applicants want those conducting these procedures to be honest, correct, sincere, and believable.

Of the three explanation rules, AI selection is inherently at risk of violating the "selection information" rule. Once again, AI is often unable to justify its methods which are especially opaque from the perspective of applicants (Cheng & Hackett, 2021). This leaves applicants unsure of how they are being evaluated, which further harms the

applicants' perceived validity of AI selection. AI selection may be perceived positively in aspects related to feedback and honesty. AI is capable of relatively immediate feedback. Also, some applicants may believe that an AI is incapable of deception. However, this reputation may be unearned. An AI may only be as honest as it is programmed to be. AI may not yet be capable of providing consistently accurate claims. Indeed, ChatGPT's website explicitly warns that its content should be fact-checked (OpenAI, 2024).

Gilliland's (1993) interpersonal treatment grouping is divided into interpersonal effectiveness of administrator, two-way communication, and propriety of questions. Interpersonal effectiveness of administrator refers to the fact that applicants view selection as fairer when communication with the person administering the selection procedure is respectful and effective. Two-way communication refers to applicants' perceiving selection methods as fairer when they allow them to respond and interact with interviewers. Propriety of questions refers to how applicants prefer the topics discussed to be related to the job, respectful, and non-invasive.

Of the three interpersonal treatment rules, AI is at risk of violating two-way communication and propriety of questions. Due to its inductive nature, many applicants may find that the criteria on which they are being judged violate the propriety of a job interview. These qualities may include facial and voice data, or even protected class characteristics. While applicants may not be told this kind of data is being collected, applicants that do know may see the collection of such data as a violation of their privacy. Two-way communication can be difficult for AI as its ability to respond to a wide variety of questions is limited. Because of this limitation, AI may often be unable to provide clarification, accommodation, or other considerations to an applicant directly.

Overall, the problems applicants have with AI can be summarized by a perceived lack of validity and a limited capability for dynamic communication. Applicants are accustomed to interacting with other human beings, but they are not used to interacting with AI. Depending on local law, employers may or may not be required to tell employees that AI is being used (Artificial Intelligence Video Interview Act, 2020). This may lead to confusion about why certain questions are being asked that AI is unable to clarify. Applicants who find out that their voice and facial data are being collected may feel that this is an invasion of their privacy and not the propriety of a job interview. Furthermore, AI's lack of immunity to bias on protected class characteristics may not be known to all applicants but will certainly be concerning to those who do.

General Attitudes about AI

General reactions to AI technology have been the topic of previous research. For instance, Fast and Horvitz (2017) created and validated a measure of general attitude toward AI. Much of the public has a negative view of AI, sometimes stemming from a wariness of its purpose and the potential for its abuse. Gherhes (2018) found that men and those with technical expertise had higher support for AI. Zhang and Dafoe (2019) studied many variables relating to AI in their research. Like Gherhes, Zhang and Dafoe found that general support for AI was greater among male individuals and those with tech expertise. Additionally, they found that AI support was higher among the educated and those making greater than \$100,000 dollars annually.

Zhang and Dafoe (2019) also found that more Americans supported the development of AI technologies than opposed it. However, Americans did not trust all organizations equally, believing some organizations would be more likely to abuse AI.

They trusted academic establishments and the military the most and tech companies the least. Facebook (now called Meta) was found to be especially distrusted. Participants believed Facebook would be more likely to abuse AI even compared to other tech companies. This finding has important implications for the use of AI hiring technologies, as applicants' reactions to AI selection may depend on a company's prior reputation.

Applicant Reactions to Non-Traditional Selection Procedures

The research on applicant reactions to AI has only recently begun, but there are similar non-traditional selection procedures that can give a clue as to what researchers should expect to discover. For instance, Folger et al. (2022) conducted similar research on digital selection procedures, that is, procedures that rely on the use of digital communication technology such as internet and video software. Digital selection procedures are often perceived as more innovative, and Folger et al. found that selection procedures that were viewed as more innovative resulted in more positive perceptions of an organization in some cases. However, Folger et al. also found that digital selection procedures resulted in more negative perceptions overall, despite their perceived innovativeness. Many of these perceptions depended on what stage of an application these innovative techniques were used. For instance, a procedure being rated as more innovative does not appear to influence applicant perception of fairness in assessment stages of interviews but does appear to influence applicant perception in interview phases where such procedures are less commonly used. Folger et al. also pointed out that standardized assessments generally have clear goals, are closely related to job function, and have less room for inconsistency between applicants, which may explain why applicants tend to perceive fewer differences between digital and traditional procedures

during the assessment phase of an application.

Folger et al. (2022) went on to state several other reasons why an applicant may react negatively to digital selection procedures. Firstly, they contain less personal interaction compared to traditional selection (McCarthy et al., 2017). When a piece of technology is seen to diminish the quality or accuracy of communication, applicants tend to find the procedure less fair. This is consistent with Gilliland's (1993) rule of procedural justice that communication in a selection procedure should be two-way. The next reason applicants may find digital selection less fair is that these procedures are more standardized (Chapman & Webster, 2001). And finally, there are privacy concerns related to these digital selection procedures (Bauer et al., 2006; Stone et al., 2013). Some other concerns that appear to impact standardized digital assessments are that algorithms may not be free of bias (Cheng & Hackett, 2019), and the possibility of technical problems (Harris et al., 2003). Some applicants may take issue with these highly digitized interview methods because it signals to them that the organization is more interested in efficiency and cutting costs (Acikgoz et al., 2020). AI selection procedures are highly digital and fit well with the Folger et al. definition for highly digitalized procedures as they require digital technology to operate, are more standardized, and will likely affect applicants' ability to communicate in similar ways.

Prior research on applicant reactions to AI has shown that applicants' reaction to AI will vary. van Esch et al. (2021) found that whether applicants would have a positive or negative reaction to AI in the selection process depended on how novel they perceived the use of AI was and their overall level of anxiety about AI selection. This corroborates prior research about selection in that innovative selection methods were found to make

organizations more attractive (Folger et al., 2022). Additionally, it sheds more light on why this effect is not universal and points to the conclusion that some applicants will react negatively to AI overall despite whatever novelty or innovativeness it may bring to the selection process.

Some recent studies have already examined applicant reactions to AI (Hess, 2022). By applying Weiner's (1985) attribution theory, Hess found that applicants generally felt they had less control over the outcome of AI selection, but also viewed it as more stable overall. This resulted in a mixed effect on procedural justice ratings for AI selection. Interestingly, while AI being perceived as more stable had a positive impact on procedural justice ratings, it did not significantly affect organizational attractiveness. Hess also found that applicants viewed AI less positively when they were not hired, which corroborated the wealth of research on self-serving bias (Miller & Ross, 1975; Ployhart & Ryan, 1997; Schmitt et al., 2004).

The literature reviewed to this point suggests many reasons AI interviews may be perceived as less fair than traditional interviews. While there is some evidence AI may be perceived as more stable than a human interviewer, AI interviews may run afoul of procedural justice rules that pertain to communication. AI may operate in ways applicants find perplexing. They may ask questions that do not plainly and apparently pertain to the job, may communicate in ways that appear improper, or may be perceived as unsuitably rigid. Finally, the very decision to use AI interviews may indicate to applicants that an organization prioritizes automation and efficiency over pursuing a mutual dialogue with applicants.

Hypothesis 1: procedural justice ratings will differ between traditional interviews and automated interviews.

- a. Consistency ratings will be higher for automated interviews compared to traditional interviews.
- b. Reconsideration opportunity ratings will be higher for traditional interviews compared to automated interviews.
- c. Feedback ratings will be higher for traditional interviews compared to automated interviews.
- d. Selection information ratings will be higher for traditional interviews compared to automated interviews.
- e. Honesty ratings will be higher for traditional interviews compared to automated interviews.
- f. Interpersonal treatment ratings will be higher for traditional interviews compared to automated interviews.
- g. Two-way communication ratings will be higher for traditional interviews compared to automated interviews.

Self-Serving Bias

Sometimes the reasons an applicant may perceive that a procedure is more or less fair may have more to do with whether the outcome benefits them. When a procedure produces an outcome unfavorable to an applicant, that applicant is likely to perceive that procedure as less fair in order to protect their self-esteem. This idea is referred to as self-serving bias (Miller & Ross, 1975; Ryan & Ployhart, 2000). The more a procedure is seen as procedurally fair, the greater the effect on an applicant's self-esteem that procedure has

(Schroth & Pradhan Shah, 2000). Self-serving bias has important implications for AI selection because applicants who feel they are disadvantaged by these procedures will perceive them as less fair. Chan et al. (1998) found that the perceived fairness of a test was correlated both with its perceived job relevance and participants' own perceived performance. This indicates that participants care both about a test's job relevance when determining how fair that test is and maintain a self-serving bias when determining job-relevance itself. In other words, applicants rate fairness as higher when job-relevance is higher, and rate job-relevance as higher when they think they performed well.

The conclusions of Chan et al. (1998) were later replicated by Schmitt et al. (2004). Applicants even exhibit self-serving bias when their performance is reviewed by a robot, blaming the robot for its negative appraisals, and taking credit for its positive appraisals (You et al., 2011). Because self-serving bias is shown to be robust across many situations it should be expected during AI selection.

Hypothesis 2: procedural justice ratings will be higher when an applicant receives a favorable outcome (hired) compared to when an applicant receives an unfavorable outcome (not hired)

Interview Self-Efficacy

One aspect of a procedurally just hiring process is that it allows applicants the chance to perform. However, self-serving bias would indicate that applicants seek the chance to perform on characteristics in which they know they excel. Therefore, a characteristic such as interview self-efficacy (Petruzziello et al., 2022) may moderate how applicants perceive the procedural justice of different selection procedures. If applicants believe they are skilled in interviews, they should prefer procedures that best

allow them to demonstrate those skills. Interview self-efficacy is an important consideration because if a selection procedure alienates those with higher interview-self efficacy, organizations may soon find that they are driving away candidates that have higher interview self-efficacy, including those who would have been high-quality hires. If an organization is alienating a group, and there are high-quality candidates among that group, then that organization will be losing access to some high-quality candidates. There may be no need for organizations to unnecessarily narrow their field of candidates in this way. Additionally, Petruzziello et al. (2022) found that individuals with high interview self-efficacy tended to also be more skilled in interviews. Thus, alienating those with high interview self-efficacy may drive away applicants with skills that may be necessary for certain jobs, such as managers, salespeople, and mediators.

Interview self-efficacy (I-SE) is a type of self-efficacy pertaining specifically to how confident a person is in job interviews. Tay et al. (2006) created a single dimension I-SE scale. They found that I-SE predicted success in job interviews and was more malleable than traits such as extraversion or conscientiousness, which also predicted interview success. Tay et al. found that successful interviews would increase a person's interview self-efficacy, and that this relationship was especially strong when that person felt like their success was internally attributable, rather than being due to external factors.

Petruzziello et al. (2022) also developed a scale of interview self-efficacy called the Multi-dimensional Job Interview Self Efficacy Scale (MJISE). However, unlike Tay et al. (2006) Petruzziello et al. developed their scale using a multidimensional approach to measure self-efficacy on five characteristics that were shown to predict success in interviews. Petruzziello et al. conducted a confirmatory factor-analysis to determine the

best fitting model for measuring interview self-efficacy. Among the single-factor, correlational, hierarchical, and bi-factorial models that were tested, Petruzzello et al. found that a bi-factorial model with five group factors and one general factor was the best fitting. The bi-factor model also had the lowest Akaike's information criterion and the lowest Bayesian information criterion values. Additionally, Petruzzello reported that all factor loadings were significant. Altogether, this indicates that the best fitted bi-factor model contained distinct dimensions that loaded simultaneously on their own respective group factor and the single general factor. The Petruzzello et al. self-efficacy dimensions were self-promotion, interaction and probing, anxiety management, logistical, and interview preparation. Petruzzello et al. found that the overall MJISE scale correlated with the Tay et al. interview self-efficacy scale with correlations of .7 or higher with self-promotion Self-Efficacy subscale correlating most highly with the Tay et al. scale. They also found that MJISE dimensions were distinct from four of the big 5 traits (Costa & McCrae, 1992), but they found that the overall scale and sub-dimensions of MJISE were moderately correlated with emotional stability, presenting mixed evidence for the scale's discriminant validity. Petruzzello et al., stated that the individual dimensions of the MJISE can be used to test for interactions with other factors.

This study will attempt to determine if three of these dimensions of self-efficacy moderate the relationship between interview method and procedural justice ratings. Self-promotion Self-Efficacy (self-promotion SE) is defined as a person's belief about their ability to present themselves as suitable for a job and to engage in impression management behaviors. interaction and probing self-efficacy (interaction SE) is defined as a person's belief that they can communicate effectively with an interviewer. Anxiety

management self-efficacy (anxiety management SE) describes a person's belief about their ability to cope with stressful situations and unpleasant emotions that can negatively affect performance in interviews. Interview preparation self-efficacy (preparation SE) is defined as a person's belief that they can effectively prepare themselves for an interview by doing things like researching the organization or rehearsing. Logistical SE is defined as a person's ability to accomplish the act of attending the interview, such as locating the venue and arriving on time.

It stands to reason that people with high self-efficacy in traditional interviews would perceive traditional interviews as being fairer. This can be explained by prior research on self-serving bias and chance to perform. Research on self-serving bias shows that people rate procedural justice higher when they perform well (Schmitt et al., 2004). Conversely, self-serving bias would indicate that those with lower interview self-efficacy would see traditional interviews as less fair because this procedure does not inherently benefit them compared to applicants who have higher interview self-efficacy. Altogether this implies a positive relationship between the dimensions of interview self-efficacy and procedural justice.

However, it may not be the case that interview self-efficacy will result in higher procedural justice ratings in AI interviews. Applicants with higher interview self-efficacy should rate procedural justice in traditional interviews higher than in AI interviews because traditional interviews both benefit them more than AI selection and provide an abundance of opportunity to demonstrate their skill in interviewing. In AI interviews, such applicants may feel they were deprived of their chance to perform. Conversely, those low in interview self-efficacy may not view AI interviews as negatively because

they do not benefit from traditional interviews compared to those with high interview self-efficacy. It should be acknowledged that AI interviews are not wholly different from traditional interviews. Therefore, when comparing the two procedures some dimensions of interview self-efficacy can be expected to have similar impacts on procedural justice ratings. Some dimensions of interview ability would seem to be almost equally relevant between AI and traditional selection. For instance, interview preparation is a relevant component of both AI and traditional selection. However, the relevance of other dimensions is more apparent for traditional interviews. For instance, the ability to engage in self-promotion is at least somewhat dependent on the presence of a human interviewer.

Although it might be possible to sway an AI (for example, by utilizing certain words or phrases that an AI might weigh more heavily), most applicants will likely not find influencing an AI as intuitive or familiar as influencing a human interviewer. It is possible that those with high self-promotion SE see their chance to influence an interviewer as a part of their chance to perform, and when denied the chance to interact with a human being, they will feel that they have been deprived of a chance to perform.

Additionally, research on self-serving bias indicates that those who are disadvantaged by a particular selection procedure will view it as less fair (Ryan & Ployhart, 2000). Because applicants with higher self-promotion SE believe that they benefit from traditional selection due to their ability to self-promote, it is hypothesized that they will rate AI selection as being lower in procedural justice compared to those with lower self-promotion SE. However, it is not clear if this effect will overcome the general tendency for applicants to view AI as less fair overall. It is possible that applicants with especially low self-promotion SE may still prefer traditional interviews,

albeit to a lesser degree than those with high self-promotion SE (shown in Figure 1).

Alternatively, applicants may even prefer AI interviews over traditional interviews if they feel that traditional interviews disadvantage them enough. This may result in a cross interaction (as shown in Figure 2). It is also unclear whether the relationship between self-efficacy and procedural justice will be positive for AI, as is hypothesized for traditional interviews. It is hypothesized that this relationship will at least be weaker for AI compared to traditional interviews (as in Figure 1), if not neutral or negative (as in Figure 2).

Hypothesis 3a: Hiring mode will moderate the relationship between self-promotion SE and procedural justice ratings so that self-promotion SE's relationship with procedural justice ratings will be positive in traditional interviews but weaker, or even negative, for AI interviews.

Hypothesis 4a: Hiring mode will moderate the relationship between self-promotion SE and chance to perform ratings, so that self-promotion SE's relationship with chance to perform ratings will be positive for traditional interviews but weaker, or even negative, for AI interviews.

Interview anxiety would also seem to be more relevant in traditional interviews. While being anxious during any selection process would be a disadvantage (Powell et al., 2018), one could also speculate that appearing anxious to a human interviewer is an even greater disadvantage. Applicants are likely going to be more self-conscious with human beings than they would when interacting with an AI interviewer, and those with higher anxiety management SE may recognize this as an advantage for traditional interviews.

Hypothesis 3b: Hiring mode will moderate the relationship between anxiety

management SE and procedural justice ratings, so that anxiety management's relationship with procedural justice ratings will be positive for traditional interviews but weaker, or even negative, for AI interviews.

Hypothesis 4b: Hiring mode will moderate the relationship between anxiety management SE and chance to perform ratings, so that self-promotion SE's relationship with chance to perform ratings will be positive for traditional interviews but weaker, or even negative, for AI interviews.

Those with high interaction SE will probably rate AI interviews as less fair. Interaction and probing questions are perhaps less relevant to AI interviewing due to current limitations in AI technology; currently AI interview methods are not known for dynamic interaction or asking probing questions. However, an applicant's ability to handle interaction and probing questions should remain highly relevant to traditional interviews. The present study will take some license with AI capabilities to present the AI interviewer as if it is currently capable of probing questions. AIs are already capable of reading and evaluating interview transcripts, and some sources have stated that chatbots were already in use for job interviews since 2019 (Joshi, 2019; Dilmeganni, 2022b). Therefore, even if it is not currently possible for an AI to ask probing questions or dynamically interact with applicants, it is not a stretch to assume that this capability will soon exist. Regardless, it is unlikely that an applicant that is told they are being interviewed by an AI would expect probing questions, therefore applicants would think their ability to answer such questions would be less relevant in such an interview. This expectation should lead to a similar moderating effect as described for self-promotion SE and anxiety management SE. Applicants with high interaction SE should see the lack of

relevance of this ability in AI interviews as a loss of advantage and a loss of chance to perform, leading to lower procedural justice ratings.

Hypothesis 3c: Hiring mode will moderate the relationship between interaction SE and procedural justice ratings so that the relationship between interaction SE and procedural justice ratings will be positive in traditional interviews and weaker, or even negative, in AI interviews.

Hypothesis 4c: Hiring mode will moderate the relationship between interaction SE and chance to perform so that the relationship between interaction SE and chance to perform will be positive for traditional interviews but weaker, or even negative, for AI interviews.

Preparation SE's effect as a moderator is difficult to predict. Coaching has been shown to be an important predictor of performance in traditional interviews (Tross & Maurer, 2008; Huffcut, et al, 2011). However, I have found no examples of AI interview coaching to date. Because of the lack of readily available AI preparation techniques, the impact of preparing for an AI interview is uncertain. Likewise, the impact of preparation SE on procedural justice ratings in AI interviews is uncertain. Therefore, the present study will make no specific hypothesis as to the moderated effects of preparation SE.

AI interviews should have different logistical hurdles than traditional interviews. AI interviews can be conducted online and can be scheduled more flexibly. The greatest logistical hurdle for AI interviews is, perhaps, the ownership of the technology required to conduct such an interview from home. Pew Research Center (2021) estimated that roughly 85% of Americans own smartphones, while around 75% own either a desktop or laptop computer. Petruzzello et al. (2022) provided items for logistical SE that were

intended to measure a person's self-efficacy in dealing with the logistical challenges inherent in traditional in-person interviews. These items may not accurately reflect an applicant's self-efficacy in dealing with logistical challenges regarding AI interviews. Therefore, no specific hypothesis will be made as to the moderated effects of logistical SE.

CHAPTER II

METHOD

Participants

One-hundred and thirty-nine participants were recruited via Cleveland State's student research pool. Only 124 participants were used in analysis. Of the 124 participants, 82 participants were female. Two participants indicated they were neither male nor female gender or indicated that they prefer not to disclose their gender. Ethnicity was measured so that the categories were mutually inclusive; participants were allowed to select more than one ethnicity. The most indicated ethnicity was White at 83, and 21 indicated they were African American/Black. Thirteen participants identified as Hispanic, 11 identified Arab/middle eastern, other ethnicities numbered less than 10. All participants were 18 years or older and the median age was 19. Age was highly skewed (2.002) and extremely kurtotic (6.150); 15 participants were older than 21. No participants were older than 40. All but two participants indicated that they expected to interview for a job position sometime in the future and within a mean of 1.26 years. Nine participants worked full-time and 81 worked part-time. Employed participants were asked to approximate the number of hours they worked a week, and the result was highly skewed (1.026) and kurtotic (1.930). The median number of approximate hours worked

per week was 20. Single participants were the most common (77) and 29 were in a cohabiting relationship. Four participants indicated that they had children. Most participants, 34, indicated that their household earned less than \$10,000 a year. 109 participants indicated they had either a high-school degree or some college. However, because all participants were required to be enrolled in a psychology course to participate, the response options “Graduated high school” and “Some college” were unlikely to be meaningfully distinct and were more likely due to different interpretations of the response options. Only 15 participants had two- or four-year degrees. No participants had graduate degrees. Five participants indicated that they were interviewed by an AI before. Most participants indicated that they had used AI before (45), but no participants considered themselves experts in AI technology. Slightly fewer participants (43) indicated that they were somewhat familiar with AI and its uses but had never personally used it before.

The required sample size needed for the study was calculated prior to data collection. Hess (2022) reported means for process justice in AI conditions and hiring manager conditions. These means were used to calculate a Cohen's d of 1.02. This Cohen's d was converted to an f^2 of .258. When estimating an interaction effect size from a main effect, Baranger (2019) suggests using four times the sample size required for the main effect to detect an interaction half the size of the main effect. A power analysis was conducted in G*power, which indicated that with 2 tested predictors and between 3 and 9 total predictors, a sample size of 47 would be required to detect Hess's effect size with an alpha of 0.05 and a power of 0.8. When multiplying this sample size by 4, according to Baranger's rule, the required sample size to detect an interaction half the effect size of

Hess's reported effect would be $n = 188$. Alternatively, assuming the interaction should have a moderate effect size of $f^2 = 0.15$ with 2 tested predictors and between 2 and 27 total predictors would result in a required sample size of $n = 108$. The sample of 124 participants did not fulfill Baranger's criteria but did fulfill the latter criteria.

Procedure

Participants were directed to a Qualtrics survey via a URL. Participants started by completing an informed consent form. They then completed the Petruzzello et al. (2022) measure of interview-self efficacy. Participants were asked for some demographic information including ethnicity, nationality, gender, and age. Participants were then randomly assigned to one of four conditions. There were two selection type conditions (AI or hiring manager), and two outcome conditions (hired or rejected) in a 2×2 between-subjects design.

Participants in each condition read a vignette (see Appendix A) that described a hiring procedure in which the applicant was asked questions over video communications software. The vignette described their responses being evaluated by either an AI or a hiring manager. The vignette told the participants that the job they were applying for is one that they desire and are qualified for. Depending on the outcome condition, participants were told they were either offered the position or rejected for the position at the end of the vignette. To maintain parity between experimental conditions, only the aforementioned changes differed between experimental groups. After reading the vignette, participants were asked to rate how procedurally just they thought the process was.

Measures

The survey contained three sections: demographic information, interview self-efficacy (I-SE), and procedural justice. Demographics were measured using self-developed items (see Appendix B for questionnaire).

Interview Self-Efficacy. I-SE was measured using Petruzzello et al.'s (2022) 20-item MJISE scale. The MJISE scale asks respondents how confident they are that they could perform various tasks related to job interviews on a 5-point Likert-type scale. The items are anchored 1 (not at all) to 5 (completely). An example item in the self-promotion dimension is "Convey a professional image." An example item in the anxiety management SE dimension is "Manage the interview related anxiety." An example item in the preparation SE dimension is "Search for information about the company." An example item in the logistical SE dimension is "Find the venue of the interview." An example item in the interaction SE dimension is "Handle probing questions." Petruzzello et al. reported an overall Cronbach's Alpha of .91. Petruzzello et al. stated that the subscales can be used separately to determine what areas an interviewer may lack confidence in.

Procedural Justice. Procedural justice was measured using the Bauer et al. (2001) Selection Procedural Justice Scale (SPJS) and the four process fairness items from the Truxillo and Bauer scale (1999). The SPJS measures procedural justice using Gilliland's (1993) procedural justice rules and the subscales of the SPJS correspond to Gilliland's rules, including chance to perform. The subscales are correlated with overall procedural justice ratings. The scale asks participants how much they agree with each given statement on a 5-point Likert scale. An example item in SPJS is "I could really show my

skills and abilities through this interview.” An example item of Truxillo and Bauer’s (1999) scale as adapted for this study is, “I feel good about the way this interview works.” The items are anchored 1 (Strongly Disagree) to 5 (Strongly Agree). Bauer et al. (2001) reported that the SPJS had a Cronbach’s alpha of .87. Truxillo and Bauer (1999) used four questions to measure general process fairness of banding. The phrasing of these questions was adapted to pertain to an interview, rather than banding.

Inclusion Criteria

There were concerns about the quality of the data, prior to hypothesis testing. The survey was timed to take 10 minutes while reading closely. However, the actual median time to completion was 8.17 minutes ($M = 9.99$, $SD = 12.22$). Though no formal pilot testing was conducted, reading every item and the vignette as quickly as possible took about 4 minutes. However, this was done by the researcher who was already familiar with the survey and knew what to expect from it. Additionally, this does not consider the sincerity or accuracy of the response itself. There were 26 respondents who completed the survey in under 4 minutes. And of those 26, eight of them answered with the same option on every scale item. This suggests that speeding was prevalent in the sample, i.e. many respondents did not take an adequate amount of time to consider the content of the survey.

There were also problems with the manipulation check. After completing the procedural justice scale, respondents were asked who conducted the interview in the vignette to test whether they noticed the hiring condition manipulation. This manipulation check included responses, “An unnamed artificial intelligence”, “an unnamed hiring manager”, “an unnamed intern”, “the unnamed owner”, “there were multiple

interviewers”, and “I don’t remember.” Of these options, the most common response was “I don’t remember” with 49 out of 139 valid responses. Only 61 respondents passed the manipulation check. Because there were five response options, participants had a 1 in 5 chance of correctly answering the manipulation check even if they answered randomly.

Greszki et al. (2015) suggest that excluding participants for speeding may only help at best in reducing noise and does little to affect marginal distribution. It was decided not to exclude speeding participants unless they also failed the manipulation check. If a participant completed the survey in under 3 minutes, they would be considered speeding. Two participants who did not complete any scale questions were first removed from the sample. There were 14 participants who were designated as speeders. All but one of these speeders was excluded for failing the manipulation check. Of these excluded participants, six of them also answered with the same value on every item on at least one of the two scales. This led to a final sample of 124 participants. The hypotheses were also analyzed a second time while excluding all participants who failed the manipulation check. These results are mentioned in the discussion section.

CHAPTER III

RESULTS

Procedural justice ratings and self-efficacy scores were scale variables. Chance to perform, a dimension of SPJS, was also a scale variable. Hiring mode and decision outcome were dichotomous variables.

The internal consistency of scale variables was analyzed using Cronbach's alpha (see Table 1). See Tables 2 and 3 for minimum and maximum scale values, means, medians, skewness, kurtosis, and standard deviations of SPJS and MJISE scales. Kolmogorov-Smirnov tests indicated that neither process fairness (Truxillo & Bauer, 1999) nor any dimension from SPJS (Bauer et al., 2021) were normally distributed, $p > .05$. However, all SPJS or MJISE scale dimensions had skewness values between -1 and 1, indicating no extreme skewness. Fisher's kurtosis values were all between -1 and 1 except for the treatment dimensions of SPJS which had a kurtosis value of 1.192. Because all skewness values were within -2 and 2, the distributions can be considered symmetrical even by conservative estimates (Hair et al., 2022). Because Fisher's kurtosis values were between -7 and 7, the variables can be considered acceptably normal.

Testing Suitability for Control Variables

The relationships between demographic variables and procedural justice DVs

were examined to assess their suitability as control variables for hypothesis tests. Demographic categories chosen by fewer than 5 participants were coded as missing, except in cases where they could be conceptually combined and collapsed with another category. There were 10 outcome variables in the study: Truxillo and Bauer's (1999) process fairness, and the 9 dimensions of Bauer et al. (2001) that the study correctly measured.

Age was a scale variable representing respondents' self-reported age. Correlations showed age was not a significant predictor of any of the dependent variables on its own, $p > .05$. Therefore, age was not used as a control variable for hypothesis tests.

For relationship status, "Prefer not to say" was coded as missing. There were fewer than 5 respondents who were either engaged or married. Engaged and married categories were combined with the cohabiting variable for an n of 20. This resulted in three categories: cohabiting/engaged/married, non-cohabiting, and single. A series of 10 one-way ANOVAs were used to assess its relationship with dependent variables.

Relationship status was a significant predictor of four out of 10 outcome variables: chance to perform, job-relatedness predictive, consistency, and two-way communication. Relationship status did not significantly predict any other dependent variables, $p > .05$. Relationship status will be used as a control variable in analyses involving these outcome variables.

Bonferonni comparisons were performed for the significant relationship status ANOVAs. Participants who were cohabiting, married, or engaged ($M = 3.68$) reported the vignette description offered significantly greater chance to perform than did participants who were in a non-cohabiting relationship ($M = 2.90$), $p = .024$. Participants who were

single ($M = 3.32$) did not report significantly different chance to perform than non-cohabiting participants or participants who were cohabiting, married or engaged, $p > .05$. Participants who were cohabiting, married, or engaged ($M = 3.85$) reported the vignette had significantly higher job-relatedness predictive than did participants who were cohabiting, $p < .001$. Single participants ($M = 3.28$) reported that the application vignette had significantly higher job-relatedness predictive than did non-cohabiting participants ($M = 2.72$), $p = .028$. Single participants and participants who were either cohabiting, married, or engaged did not significantly differ in reported job-relatedness predictive, $p = .079$. Participants who were cohabiting, married, or engaged ($M = 4.22$) reported the vignette description depicted significantly higher consistency than did non-cohabiting participants ($M = 3.63$), $p = .045$. Single participants ($M = 3.74$) were not significantly different than either group in reported consistency, $p > .05$. Participants who were either cohabiting, married, or engaged ($M = 3.89$) reported the vignette description depicted significantly higher two-way communication than did non-cohabiting participants ($M = 3.21$), $p = .024$. Single participants ($M = 3.69$) reported significantly higher two-way communication compared to the non-cohabiting group ($p = .031$) but were not significantly different than the group of cohabiting, married, or engaged participants ($p = 1.00$).

Because relationship status significantly predicted chance to perform, job relatedness predictive, consistency, and two-way communication, relationship status will be used as a control variable in hypotheses that use these as dependent variables:

Hypothesis 1 and Hypothesis 4 (see Table 4 for hypothesis summaries).

For gender, “nonbinary/other gender” and “prefer not to say” were coded as

missing due to low representation in the sample. A series of t-tests showed Gender was not a significant predictor of any dependent variables, $p > .05$. Therefore, gender was not used as a control variable in hypothesis tests.

Due to the wide number of ethnicity categories and combinations, as well as low representation among some ethnicities or combinations of ethnicities, it was decided that ethnicity would be coded as two dichotomous mutually inclusive variables. The first variable was “Black,” and the second variable was “White.” At least 20 respondents identified with one of these ethnicities. For example, respondents who did not indicate they were either Black or White were coded as zero on both variables, while respondents who indicated they were both White and Black were coded as 1 on both variables. Participants who did not answer the question or selected “prefer not to answer” were coded as missing. Multiple linear regressions were used to assess the relationship between the dependent variables and these two ethnicity variables. The two mutually inclusive ethnicity variables were entered as predictors into 10 regressions for each respective outcome variable. The ethnicity variables were used as controls for each of the hypothesis tests involving a dependent variable which they predict. Ethnicity did not predict outcome variables: process fairness, chance to perform, job relatedness predictive, consistency, openness, treatment, two-way communication, or job relatedness content. Ethnicity significantly predicted information known, $F(2, 119) = 4.360, p = .015$. While controlling for whether a respondent was White, Black participants’ ratings of information known was -.483 lower than the average of participants who were neither White nor Black ($M = 4.126$), $t(121) = -2.334, p = .021$. Whether or not a respondent was White did not significantly predict *information known* while controlling for whether a

participant was Black when compared to a participant who was neither White nor Black, $b = .058$, $t(121) = .351$, $p = .726$. Ethnicity also significantly predicted propriety of questions, $F(2,116) = 3.142$, $p = .047$. While controlling for whether a respondent was White, Black participants' ratings of propriety of questions was .513 lower than the average of participants who were neither White nor Black ($M = 4.092$), $t(118) = -2.440$, $p = .016$. Whether or not a respondent was White did not significantly predict propriety of questions while controlling for whether a participant was Black, $b = -.0301$, $t(118) = -1.768$, $p = .08$.

Employment status was a dichotomous self-report measure of whether the respondent was employed. A series of t-tests showed Employment status was not a significant predictor of any dependent variables, $p > .05$. Employment status was not used as a control variable in any hypothesis tests.

Hours employed was a scale variable representing the estimated hours respondents worked at their jobs each week if they already indicated that they were employed. The relationship between hours employed and dependent variables was assessed using correlations. Hours employed did not significantly predict any dependent variables, $p > .05$. Hours employed was recoded to include those who indicated they were not employed to be listed as having worked zero hours. After this change, hours employed still did not predict any dependent variables, $p > .05$. Number of hours employed was not used as a control variable for hypothesis tests.

Participants were asked their annual household income. Because this was a self-report measure and many participants were unlikely to be able to offer exact figures while responding to the survey, income was measured on an ordinal scale. These were the

levels of income respondents could select from: Under \$10,000, \$10,001-\$20,000, \$20,001-\$30,000, \$30,001-\$40,00, \$40,001-\$50,000, \$50,001-\$60,00, \$60,001-80,000, \$80,001-\$100,00, over 100,000. A series of Spearman correlations indicated income did not significantly predict any dependent variables, $p > .05$. Income was not used as a control variable for hypothesis tests.

Education levels were measured ordinally in a self-report measure of highest degree earned. Because participants were recruited through an undergraduate research pool, all participants were enrolled in a college course. While there were some respondents who indicated that they either had completed high school and others that indicated they completed some college, all participants were required to be enrolled in a course to be eligible for the study, and so these responses essentially meant the same thing. Because of this, there was almost no variability in education levels. Education levels were not further analyzed for suitability as control variables.

Respondents were asked if they had ever been interviewed by an AI before. Only four of the 124 respondents who answered the question indicated that they had. Due to this low variability, this variable was not further analyzed for suitability as a control variable.

Respondents were asked if they had ever been interviewed for a job position before. A series of t -tests indicated that whether they had been interviewed did not significantly predict any dependent variable, $p > .05$. Therefore, this variable was not used as a control variable in further analysis. Participants were asked whether they had children. Only four participants indicated they had children. Due to low variability, whether participants had children was not analyzed any further for suitability as a control

variable in hypothesis tests.

Participants were asked how familiar they were with AI before the study. AI familiarity was measured ordinally from 1-6, one being “I’d never heard of AI before today” and 6 being “I am an expert in AI technology.” No participant picked value six, thus the highest familiarity a participant reported was value 5, “I am familiar with the technology and use it often.” A one-way ANOVA showed AI familiarity significantly predicted chance to perform ($F(4, 117) = 2.66, p = .034$), job relatedness predictive ($F(4, 118) = 3.992, p = .004$), and job relatedness content ($F(4, 118) = 2.825, p = .028$). However, a Bonferroni post-hoc found no significant differences between groups, $p > .05$. An LSD post-hoc was also performed for chance to perform and Job relatedness predictive, ($p > .05$). Participants who stated, “I am somewhat familiar with AI technology and its uses but do not use it myself” ($M = 3.5875$) rated the chance to perform significantly higher than participants who stated, “I have used AI technology before but am not an expert” ($M = 3.0341$). Participants who stated, “I’d never heard of AI technology before today” ($M = 3.318$) rated chance to perform significantly higher than those who stated, “I have used the technology before but am not an expert,” and those who stated, “I’d heard of AI technology but know very little about it. It is worth noting that the results of this post-hoc suggested non-linearity, suggesting it should not be treated ordinally. Because the ANOVA test was significant, familiarity with AI was used as a categorical control variable in Hypothesis 1 and Hypothesis 4 (see Table 4 for hypothesis summary).

Hypothesis 1

Hypothesis 1 was tested using a MANOVA (see Table 4 for hypothesis summary). The predictor was hiring mode. The dependent variables were chance to perform, job relatedness predictive, information known, consistency, openness, treatment, two-way communication, propriety of questions, and job relatedness content.

Reconsideration opportunity and feedback were both excluded. This was because of an error in data collection in which the respondents did not receive all the items for these two subscales. Ethnicity variables Black and White, as well as dummy coded relationship status variables were used as covariates because they have relationships with dependent variables in this analysis. Familiarity with AI was used as a covariate. The study's other manipulation, whether participants received a job offer in their vignette, was also used as a control variable. Levene's test found groups did not differ in variance, $p > .05$. No other dependent variable had significant differences in error variances, $p > .05$. Box's M indicated problems with equality of covariance matrices, $p < .001$. Residual plots did not indicate heteroscedasticity for any variables. Pillai's Trace was not significant for any predictors, $p > .05$. See Table 5 for Multivariate tests. Hypothesis 1 was not supported. No post-hoc tests were examined due to failure to reject the multivariate omnibus test. Therefore, sub hypotheses will not be interpreted or discussed. See Table 6 for Hypothesis 1 regressions.

Hypothesis 2

Hypothesis 2 was tested using hierarchical multiple regression (see Table 4 for hypothesis summary). The dependent variable was process fairness rating (Truxillo & Bauer, 1999). The predictor variable was hiring outcome (received job offer/ did not

receive job offer). The study's other manipulation, hiring mode, was used as a control variable. No demographic variables were found suitable for use in Hypothesis 2 as control variables. Because Hiring mode was treated as a control variable, an interaction between it and job offer was not tested.

The VIF value between the predictor (job offer) and the control variable (hiring mode) was lower than 2.5 indicating no multicollinearity issues, $VIF = 1.002$. Outliers were assessed using Mahalanobis distance. In all instances Mahalanobis distance values were assessed using a chi-square distribution with a significance level of .001 as suggested by Tabachnick and Fidell (2013) The maximum Mahalanobis distance value was 2.157. Because there were two predictors in the model, the Mahalanobis distance values were compared to a chi-square distribution with two degrees of freedom. This indicated no outliers at the .001 level, observation with lowest $p = .340$. The P-P plot indicated no issues with residual normality, and residual-predicted plots indicated no heteroscedasticity. Kolmogorov-Smirnov test found that residuals were significantly different than normal. However, skewness was less than 1 and Kurtosis was less than 2 (See Table 7).

The multiple regression found that control variable hiring mode significantly predicted process fairness on its own, $F(1, 120) = 5.434, p = .021$. Furthermore, adding the job offer predictor in Model 2 significantly improved prediction of process fairness, $\Delta R^2 = .035, F(1, 119) = .035, p = .035$. Because Model 2 was a significant improvement, following interpretations will pertain to this model. Participants who read a vignette which described a traditional job interview with no job offer had a mean process fairness of 3.751, constant $b = 3.751, t(120) = 23.976, p < .001$. Participants who read a vignette

of an AI interview rated process fairness .437 lower on average, $t(120) = -2.467, p = .015$. Participants who read a vignette that offered a job rated process fairness .377 higher on average, $t(120) = 2.130, p = .035$. These findings supported Hypothesis 2. See Table 8 for Hypothesis 2 coefficients.

Hypothesis 3

Hypothesis 3a, Hypothesis 3b, and Hypothesis 3c were each tested using multiple regression (see Table 4 for hypothesis summary). The dependent variable for all of these regressions was process fairness ratings. Because hiring mode is a dichotomous variable it was dummy coded. For each respective regression, moderation was assessed by testing the interactions between the SE variable and hiring mode.

For Hypothesis 3a the predictor variables were hiring mode and self-promotion SE. The dependent variable was process fairness. Model 1 included the aforementioned predictors and did not include their interaction. Model 1 significantly predicted process fairness, $R^2 = .104, n = 120, F(2, 117) = 6.782, p = .002$. A scatter plot indicated no sign of heteroscedasticity. The maximum Mahalanobis' distance value was 12.606. Because there were three predictors in the model, the Mahalanobis' distance values were compared to a chi-square distribution with three degrees of freedom. This indicated no outliers at the .001 level (Tabachnick & Fidell, 2013), observation with lowest $p = .006$. The standardized residuals appeared approximately normally distributed in histogram. A Kolmogorov-Smirnov test found that residuals were not normally distributed. However, skewness values were less than 1 and Kurtosis values were less than 2. See Table 7 for residuals and normality test. The P-P plot showed little sign of bowedness. The VIF between hiring mode and self-promotion SE was 1.026, indicating no issues with

multicollinearity. Both hiring mode ($b = -.351, t(119) = -1.993, p = .049$) and self-promotion ($b = .318, t(119) = 2.742, p = .007$) were significant predictors of process fairness; participants in the AI condition rated process fairness lower than the traditional interview condition, and participants with higher self-promotion self-efficacy tended to rate process fairness as higher. Model 2 included the interaction between hiring mode and self-promotion SE. However, Model 2 was not a significant improvement over Model 1, $\Delta R^2 = 0.000, F(3, 116) = .033, p = .857$. Because of this, the coefficients for Model 2 will not be discussed. See Table 9 for Hypothesis 3a coefficients.

For Hypothesis 3b the predictor variables were Hiring mode and anxiety management SE. Model 1 did not include the interaction and significantly predicted process fairness, $R^2 = .109, n = 121, F(2, 118) = 7.221, p < .001$. A scatter plot indicated no sign of heteroscedasticity. The maximum Mahalanobis' distance value was 14.395. Because there were three predictors in the model, the Mahalanobis' distance values were compared to a chi-square distribution with three degrees of freedom. Tabachnick and Fidell (2013) suggest using the .001 significance level, which indicated no outliers in this analysis, observation with lowest $p = .002$. The standardized residuals appeared slightly negatively skewed in histogram. A Kolmogorov-Smirnov test found that residuals were not normally distributed. However, Skewness values were less than 1 and Kurtosis values were less than 2 (see Table 7 for residuals). The P-P plot showed little sign of bowedness. The VIF between hiring mode and anxiety management SE was 1.006, indicating no issues with multicollinearity. Hiring mode was not a significant predictor of process fairness $b = -.362, t(118) = -2.074, p = .040$. However, anxiety management SE did predict process fairness, $b = .307, t(118) = 3.015, p = .003$; participants with higher

anxiety management tended to rate process fairness higher. Model 2 included the interaction between hiring mode and anxiety management SE, but it was not a significant improvement over Model 1, $\Delta R^2 = 0.014$, $F(1, 117) = 1.924$, $p = .186$. Because of this, Model 2 will not be discussed (see Table 10).

For Hypothesis 3c the predictor variables were Hiring mode and interaction SE. Model 1 did not include the interaction and significantly predicted process fairness, $R^2 = .078$, $n = 121$ $F(2, 118) = 5.006$, $p = .008$. A scatter plot indicated no sign of heteroscedasticity. The maximum Mahalanobis' distance value was 17.211. Because there were three predictors in the model, the Mahalanobis' distance values were compared to a chi-square distribution with three degrees of freedom. This indicated that one observation was an outlier at the .001 level (Tabachnick & Fidell, 2013), observation's $p = .0006$. This observation had a Cook's d value of .21, indicating that it had low influence, $.21 < 1$ (Cook & Weisberg, 1982). The standardized residuals appeared approximately normally distributed in histogram. A Kolmogorov-Smirnov test found that residuals were not normally distributed. However, skewness values were less than 1 and Kurtosis values were less than 2. See Table 7 for residuals and normality test. The P-P plot showed little sign of bowedness. The VIF between hiring mode and interaction SE was 1.028, indicating no issues with multicollinearity. Both hiring mode ($b = -.368$, $t(118) = -2.038$, $p = .044$) and interaction SE ($b = .207$, $t(118) = 2.049$, $p = .043$) were significant predictors of process fairness; participants in the AI condition tended to rate process fairness as lower than participants in the traditional interview condition, and participants with higher interaction SE tended to rate process fairness as higher in general. Model 2 included the interaction between hiring mode and interaction SE

However, Model 2 was not a significant improvement over Model 1, $\Delta R^2 = 0.004$, $F(1, 117) = .533$, $p = .467$. Because of this, Model 2 coefficients will not be discussed (see Table 11).

Hypothesis 4

Hypothesis 4 was tested using multiple regressions (see Table 4 for hypothesis summary). The dependent variable for all of these regressions was opportunity to perform. Because hiring mode is a dichotomous variable it was dummy coded. For Hypothesis 4a the predictors were hiring mode and self-promotion SE. Relationship status was dummy coded into two variables: Married/engaged/cohabiting, and non-cohabiting. Familiarity with AI was dummy coded into 4 variables: Use AI often, somewhat familiar, Use AI, Heard of AI. Model 1 did not include the interaction and significantly predicted process fairness, $R^2 = .194$, $n = 119$, $F(8, 110) = 3.309$, $p = .002$. A scatter plot indicated no sign of heteroscedasticity. The maximum Mahalanobis' distance value was 34.015. Because there were eight predictors in the analysis, the Mahalanobis' distance values were compared to a chi-square distribution with eight degrees of freedom. This indicated two outliers at the .001 level (Tabachnick & Fidell, 2013). The highest Mahalanobis d observation = 34.01536, $p = .00004$, had a Cook's d of .18440. The second highest Mahalanobis d value = 26.80683, $p = .00076$, had a Cook's d value of .00625. This indicates that both outliers had low influence, Cook's $d < 1$. The standardized residuals appeared approximately normally distributed in histogram. A Kolmogorov-Smirnov test found that residuals were not significantly different than normal (see Table 7 for residuals). The P-P plot showed little sign of bowedness. The VIF values for all variables were less than 4. Because no variables had VIF higher than

five, this indicated no issues with multicollinearity. Hiring mode was not a significant predictor of chance to perform while controlling for self-promotion SE, relationship status, and AI familiarity, $b = -.324$, $t(110) = -1.906$, $p = .059$. Self-promotion was a significant predictor of chance to perform when controlling for hiring mode, relationship status, and familiarity with AI. Participants who had higher self-promotion SE reported having more chance to perform, $b = .272$, $t(110) = 2.335$, $p = .021$. Relationship was not a significant predictor of chance to perform while controlling for self-promotion and hiring mode, $p > .05$. Participants who indicated that they had used AI before but were not experts reported they had significantly less chance to perform compared to participants who had never heard of AI, $b = -.707$, $t(110) = -2.223$, $p = .028$. Model 2 included the interaction between hiring mode and self-promotion SE. However, Model 2 was not a significant improvement over Model 1, $\Delta R^2 = 0.007$, $F(1, 109) = .985$, $p = .323$. Because of this, Model 2 coefficients will not be discussed (see Table 12).

For Hypothesis 4b the predictors were hiring mode, relationship status, and anxiety management SE. Model 1 did not include the interaction and significantly predicted process fairness, $R^2 = .249$, $n = 120$ $F(8, 111) = 4.596$, $p < .001$. A scatter plot indicated no sign of heteroscedasticity. Because there were eight predictors in the analysis, the Mahalanobis' distance values were compared to a chi-square distribution with eight degrees of freedom. This indicated that there were three outliers at the .001 level (Tabachnick & Fidell, 2013). The highest Mahalanobis d value = 34.23350, $p = .00004$. The second highest Mahalanobis d value = 29.91034, $p = .00022$. The third highest Mahalanobis d value = 26.55824, $p = .00084$. All observations in this model had Cook's d values of less than 1, indicating low influence. The standardized residuals

appeared approximately normally distributed. A Kolmogorov-Smirnov test found that residuals were not significantly different than normal (see Table 7). The P-P plot showed little sign of bowedness. The VIF values for all variables were less than 4. Because no variables had VIF higher than five, this indicated no issues with multicollinearity. Both hiring mode ($b = -.381, t(111) = -2.332, p = .022$.) and anxiety management ($b = .338, t(111) = 3.537, p < .001$) were significant predictors of chance to perform. Relationship status did not predict chance to perform while controlling for anxiety management and hiring mode, $p > .05$. Those who stated they had used AI before but were not experts indicated they had significantly less chance to perform compared to those who had never heard of AI, $b = .789, t(111) = -2.543, p = .012$. Those who stated they had heard of AI but knew very little about it also reported they had less chance to perform, $b = .627, t(111) = -2.021, p = .046$. Model 2 included the interaction between hiring mode and anxiety management SE. However, Model 2 was not a significant improvement over Model 1, $\Delta R^2 = 0.006, F(1, 110) = .914, p = .341$. Because of this, Model 2 coefficients will not be discussed (see Table 13).

For Hypothesis 4c, the predictor variables were hiring mode and interaction SE. For each respective regression, moderation was assessed by testing the interactions between the SE variable and hiring mode. Model 1 did not include the interaction between hiring mode and interaction SE. Model 1 significantly predicted chance to perform, $R^2 = .230, n = 120, F(8, 111) = 4.152, p < .001$. A scatter plot indicated no sign of heteroscedasticity. Because there were eight predictors in the analysis, the Mahalanobis' distance values were compared to a chi-square distribution with eight degrees of freedom. This indicated that there were two outliers at the .001 level

(Tabachnick & Fidell, 2013), Highest Mahalanobis d value = 34.12694, $p = .00004$. The second highest Mahalanobis d value = 26.25570, $p = .00095$. All observations had Cook's d values less than 1, indicating they had low influence. The standardized residuals appeared approximately normally distributed in histogram (see Table 7 for residuals and normality test). A Kolmogorov-Smirnov test found that residuals were not significantly different than normal. The P-P plot showed little sign of bowedness. The VIF values for all variables were less than 4. Because no variables had VIF higher than five, this indicated no issues with multicollinearity. Hiring mode did not significantly predict chance to perform while controlling for other variables in model 1, ($b = -.325$, $t(111) = -1.953$, $p = .053$). Interaction SE did predict chance to perform ($b = .301$, $t(111) = 3.232$, $p = .002$) were significant predictors of chance to perform. Relationship status did not significantly predict chance to perform while controlling for hiring mode and interaction SE, $p > .05$. Familiarity with AI predicted chance to perform in model 1, $p < .05$ (see Table 14). Model 2 included the interaction between hiring mode and interaction SE. However Model 2 was not a significant improvement over Model 1, $\Delta R^2 = 0.002$, $F(1, 110) = .281$, $p = .597$. Because of this, Model 2 coefficients will not be discussed (see Table 14).

Chapter IV

DISCUSSION

The study has supported findings in prior literature but failed to find evidence of a connection between interview self-efficacy and perception of AI interviews. Prior literature about AI usage in interviews being generally negative was partially supported. No evidence that participants perceived AI as less fair was found in Hypothesis 1's analysis for any dimension of PJ. Hypothesis 1 results were unable to identify what dimensions of PJ were most affected by use of AI interviews. In Hypothesis tests for 2, 3, and 4, participants who read the vignette of an AI interview rated chance to perform, and process fairness in general, lower than participants who read the traditional interview vignette. Unlike Hypothesis 1, these tests used a direct measure of PJ (Truxillo & Bauer, 1999). This partially replicates the prior research on general distrust of AI selection methods. As of late 2023, the public perception around AI continues to change as it becomes a more mainstream topic. Within this study's creation alone, both Dall-E and ChatGPT were released, and AI went from a niche topic to a common headline. Future research may naturally find that perceptions of AI interviews continue to change as the public becomes more aware, more educated, and/or the tools being used change.

In retrospect, Hypothesis 1 could have included a Hypothesis regarding chance to

perform. It could be reasoned that if participants may feel generally uncomfortable with AI selection and generally rate it as less fair (Hess, 2022), then this could be partly due to a perception that they had less opportunity to demonstrate their abilities. While the study did not make this hypothesis at the onset, it may be noted that the Hypothesis 1 analysis did not find that hiring condition or a job offer had any effect on chance to perform or any other dimension of SPJS (Bauer et al., 2001). However, Hypothesis 4 analyses did find that hiring condition had an effect on opportunity to perform.

Self-Serving Bias

Prior research on self-serving bias was successfully replicated, but this study's hypotheses regarding interview self-efficacy were not supported. Participants who read a vignette in which they were told they received a job offer rated process fairness as significantly higher than those who read a vignette that did not include a job offer, supporting Hypothesis 2. Additionally, this study extended research on self-serving bias by finding that participants rated interviews as significantly fairer when their interview self-efficacy was high, demonstrating that participants may view such hypothetical interviews more favorably when they believe they would fare well in them. However, this study hypothesized that there would be an interaction between interview self-efficacy and hiring conditions such that the higher a person's interview self-efficacy was, the fairer they would rate the traditional interview in particular. It was suggested that a person with higher interview self-efficacy would feel better advantaged by a traditional interview compared to an AI interview and may even feel deprived of their chance to perform in an AI interview. This did not appear to be the case in this study; the relationship between self-efficacy and fairness ratings was not significantly affected by the type of interview.

This was true for both general fairness ratings as well as chance to perform, more specifically. Nor did it matter what dimension of interview self-efficacy was being considered. Therefore, Hypotheses 3 and 4 were not supported. The findings that interview self-efficacy is related to perceived fairness of interviews aligns well with prior research on self-serving bias and seem logical in retrospect. However, these main effects were not hypothesized, and future research may wish to replicate them.

The Effects of MJISE on Procedural Justice

There are multiple possible explanations for the findings regarding Hypotheses 3 and 4. The first, of course, is that this interaction does not exist. It may be that applicants with higher interview self-efficacy do not feel especially disadvantaged by AI interviews, and that applicants with lower interview self-efficacy have no particular tolerance or preference for AI interviews. It is possible that this self-serving bias could still emerge as the public becomes more aware of AI selection methods and how they might affect their outcome, but this emergence seems most likely under the premise that participants skilled in interviews are actually disadvantaged by AI interviews. Petruzziello et al. (2022) found that interview self-efficacy's general factor positively correlates to skill in interviews, but this study did not directly measure participants' skill in interviews. It may seem an intuitive conclusion, but whether applicants with lower or higher interview skills are disadvantaged or advantaged in either type of interview has not been tested as of this writing. Future research may consider testing this. Beyond the implications for the type of person that an AI interview may be selecting for, the existence of a real advantage or disadvantage in an AI interview would further imply the possibility of a self-serving bias.

It is possible that self-efficacy in traditional interviews and self-efficacy in AI

interviews are not distinct constructs. There is also the possibility that interview self-efficacy is related to a new unestablished construct. Bandura (2006) suggests that self-efficacy is task specific and not a general trait, so individuals with high interview self-efficacy having high self-efficacy in all other tasks seems an unlikely explanation. However, Chen et al. (2001) state that stable traits such as conscientiousness and cognitive ability can predict general self-efficacy. Still, interview self-efficacy may be moderately related to distinct constructs of self-efficacy in AI related tasks. To this point, no study has examined individuals' self-efficacy in AI prompt engineering, for instance. That is, it may be that those who believe they are skilled in interviews also believe they can better interact with AI in other ways or that types of AI interaction exist as distinct self-efficacies. Future research may wish to explore more about AI related skills, tasks, and their corresponding self-efficacies.

The next explanation for results from Hypotheses 3 and 4 is that the study was not powerful enough to detect the effect. Most participants failed the manipulation check, suggesting that the strength of the manipulation was low, and that the study's fidelity may have also been low. A closer to real life situation may produce a stronger effect that a study with greater power may be able to detect, if the effect does exist.

Limitations

The exclusion criteria were not formally tested for validity. Participants were excluded if they did not answer sufficient items to be analyzed, or if they both were designated as speeders (took less than three minutes) and failed the manipulation check. It is possible that these criteria reduced the study's power enough to alter the conclusions, and it may also be possible that stricter criteria could have made the sample more

representative. However, the analyses were conducted again with all participants and results were not substantially different than what has been discussed. Because no pilot testing was performed prior to the study, no conclusions will be made as to which is more likely, but it remains a limitation worth considering. Future research may wish to more carefully sample with attrition in mind.

Only 61 participants passed the manipulation check for hiring condition. The study chose to exclude participants that failed the manipulation check only if they also completed the survey in less than three minutes. However, this calls into question the strength of the manipulation. The present study's use of vignettes may explain the poor strength of the manipulation and may also suggest that the study had low fidelity; the study lacked realistic conditions. This may present external validity issues, as the experience of reading a vignette may not generalize to real interviews. Future research may consider better simulating or using actual interviews to both increase the strength of the manipulation and improve fidelity.

The entirety of the Bauer et al. (2001) SPJS scale was not included due to an error in collection. Instead of containing all of the items from the reconsideration and feedback dimensions of SPJS, several questions were repeated in their place. This meant that the reconsideration and feedback dimensions of the SPJS scale could not be scored, and the study could not test hypotheses H1a and H1b. The study was also not able to test other hypotheses using overall procedural justice ratings using SPJS as planned and could only test Hypotheses 2 and 3 using Truxillo and Bauer's (1999) direct measure of procedural fairness. Ideally, PJ would have been measured in a consistent way across all tests. Other than reconsideration and feedback dimensions, all other parts of the SPJS scale were

intact and used as planned in Hypothesis 1 and 4. Future research may consider following up on this by testing the full Bauer et al. (2001) PJ scale. The condition in which participants read about a hiring manager and read a job offer had a duplicate vignette; in the “hiring manager rejected” condition the vignette was presented, then presented again. This may have given the participants an impression of sloppiness that could have influenced their PJ ratings. However, the hiring manager condition (the traditional interview) still had higher fairness scores, indicating that this was likely not a problem overall.

There are other important limitations to consider regarding this study’s generalizability. As previously stated, the study had low fidelity and real-world situations may differ in ways that were not represented in the ratings of this study’s vignettes. It is expected that in a real situation the effects found in this study would be stronger. The study had a low sample size, and many participants failed the manipulation check. All of these factors likely reduced the power or effect size of the study and may explain why some hypotheses could not be confirmed. Additionally, all participants were students in the same region of the United States at the same university during the same semester. Student ages were not widely distributed, and they had nearly identical education levels. This study represents only a very narrow sample of what one could expect to find in the real job market. Future research may wish to sample from a more diverse population to better test the external validity of these conclusions, as the effects may be stronger or weaker in other populations and settings.

Future Research: Connecting Results to Other Frameworks

Self-serving bias is a theory that is heavily based in the framework of Weiner’s

(1985) attribution theory (Ployhart & Harold, 2004). Self-serving bias is an attributional error where individuals tend to prefer attributing success to their disposition while attributing their failure to situations. Ployhart and Harold provide a framework for how applicant reactions function in the context of attribution theory. While other studies have included mediators from these models (Hess, 2022), this study did not incorporate many elements of it. Future research may wish to further investigate whether interview self-efficacy's effects on perceived fairness of interviews can be explained using attribution theory by incorporating variables from the theory such as controllability, causality, stability. For instance, Hess found that participants whose outcome was determined by an AI had higher perceptions of stability for the application process, though this did not fully compensate for other negative perceptions of AI. Future research may wish to study whether MJISE dimensions have relationships with a preference for stable hiring procedures.

Applicant reactions research is concerned with how an application process affects the applicant's perceptions of the organization. No direct data were collected about how applicants may react to an organization that used the hiring method in this study, nor how interview self-efficacy could affect views about that organization. To further connect these findings to the rest of applicant reactions literature, future research may consider asking these questions to investigate the practical implications of these effects for organizations.

Practical Implications

Organizations should be aware that the types of application processes they use do influence applicant's opinions of the organization (Kohn & Dipboye, 1998; Latham &

Finnegan, 1993). Participants in this study had a generally negative reaction to the use of AI in this interview and rated the interview as less fair. Furthermore, participants who had higher interview self-efficacy were likely to rate the application process as less fair overall. This suggests that the use of AI may be alienating applicants, and that interviews in general are likely to alienate any participants who feel they do not fare well in interviews. This has implications for organizations who are not seeking to hire applicants for their interview ability, as an interview is likely to mistakenly alienate qualified applicants. However, finding a suitable substitute for interviews that applicants with lower interview self-efficacy will rate as more just may prove challenging, as many alternatives to traditional interviews are digital and asynchronous which have their own applicant reaction issues (Folger et al., 2022).

AI may provide many operational benefits to organizations who use them, but organizations should be aware of its effect on applicant perceptions. As in prior research, self-serving bias indicates that organizations in general should be wary of using application methods that their desired imagined employee would not fare well in, as it may alienate the very employee that the organization is seeking. Organizations who are concerned about this may wish to avoid AI application processes. Alternatively, organizations may wish to find ways of accommodating the effects of applications methods that their participants may find worrisome. Gilliland's (1993) procedural justice rules should provide useful guidelines for what elements should be addressed in modification. As an example, one such modification may be to include ways for AI to provide effective and immediate feedback to participants, preserving the automated benefits of AI while addressing its issues with Gilliland's feedback dimension of

procedural justice. This capability does not seem far off, as during the writing and conducting of this study, language processing models such as ChatGPT were released, though it remains to be seen how effective such feedback may be. Future research may wish to study in greater detail how the characteristics of the applicant affect the individual's need for each of Gilliland's dimensions, and what kind of modifications to AI selection may improve these justice perceptions.

REFERENCES

- Acikgoz, Y., Davison, K. H., Compagnone, M., & Laske, M. (2020). Justice perceptions of artificial intelligence in selection. *International Journal of Selection and Assessment*, 28(4), 399-416. <https://doi.org/10.1111/ijsa.12306>
- Asfahani, A. M. (2022). The Impact of artificial intelligence on industrial-organizational psychology: A systematic review. *Journal of Behavioral Science*, 17(3), 125-139.
- Bandura, A. (2006). Guide for constructing self-efficacy scales. Self-efficacy beliefs of adolescents. In T. Urdan, & F. Pajares (Eds.), *Self-efficacy beliefs of adolescents* (pp. 307–337). Charlotte, North Carolina: Information Age Publishing.
- Baranger, D. (2019, August 6). *Interaction analyses- how large a sample do I need?* (Part 3). Davidbaranger.com
- Bauer, T. N., Truxillo, D. M., Sanchez, R. J., Craig, J. M., Ferrara, P., & Campion, M. A. (2001). Applicant reactions to selection: Development of the selection procedural justice scale (SPJS). *Personnel Psychology*, 54(2), 388–420. <https://doi.org/10.1111/j.1744-6570.2001.tb00097.x>
- Brooks, R. (2018, April 27). *The Origins of “Artificial Intelligence.” Rodney Brooks. Robots, AI, and other Stuff.* <https://rodneybrooks.com/forai-the-origins-of-artificial-intelligence/>
- Chen, G., Casper, W. J., & Cortina, J. M. (2001). The roles of self-efficacy and task complexity in the relationships among cognitive ability, conscientiousness, and work-related performance: A meta-analytic examination. *Human Performance*, 14(3), 209–230. https://doi.org/10.1207/S15327043HUP1403_1
- Chapman, D. S., Uggerslev, K. L., Carroll, S. A., Piasentin, K. A., & Jones, D. A. (2005).

- Applicant attraction to organizations and job choice: a meta-analytic review of the correlates of recruiting outcomes. *Journal of Applied Psychology*, 90(5), 928–944. <https://doi.org/10.1037/0021-9010.90.5.928>
- Chapman, D. S., & Webster, J. (2001). Rater correction processes in applicant selection using videoconference technology: The role of attributions. *Journal of Applied Social Psychology*, 31(12), 2518–2537. <https://doi.org/10.1111/j.1559-1816.2001.tb00188.x>
- Charmarro, T. & Ahtktar, R. (2019). *Should companies use AI to assess job candidates?*. Harvard Business Review. <https://hbr.org/2019/05/should-companies-use-ai-to-assess-job-candidates>
- Chan, D., Schmitt, N., Jennings, D., Clause, C. S., & Delbridge, K. (1998). Applicant perceptions of test fairness: Integrating Justice and self-serving bias perspectives. *International Journal of Selection and Assessment*, 6(4), 232–239. <https://doi.org/10.1111/1468-2389.00094>
- Cheng, M. M., & Hackett, R. D. (2019). A critical review of algorithms in HRM: Definition, theory, and practice. *Human Resource Management Review*, 31(1), Article 100698. <https://doi.org/10.1016/j.hrmr.2019.100698>
- Costa, P. T., & McCrae, R. R. (1992). The five-factor model of personality and its relevance to personality disorders. *Journal of Personality Disorders*, 6(4), 343–359. <https://doi.org/10.1521/pedi.1992.6.4.343>
- Cook, R. D., & Weisberg, S. (1982). *Residuals and influence in regression*. New York: Chapman and Hall.
- Cropanzano, R., Bowen, D. E., & Gilliland, S. W. (2007). The management of

- organizational justice. *Academy of Management Perspectives*, 21(4), 24-48.
<https://doi.org/10.5465/AMP.2007.27895338>
- Cropanzano, R., & Wright, T. A. (2003). Procedural justice and organizational staffing: A tale of two paradigms. *Human Resource Management Review*, 13(1), 7-39.
[https://doi.org/10.1016/S1053-4822\(02\)00097-9](https://doi.org/10.1016/S1053-4822(02)00097-9)
- Dilmeganni, C. (First published 2020, Updated 2022) Recruiting AI in 2022: Guide to augmenting the hiring team. Site: AI Multiple. Accessed on March 22 2022.
<https://research.aimultiple.com/recruiting-chatbot/>
- Dilmeganni, C. (First published 2020, Updated 2022) Recruiting AI in 2022: In-Depth Guide into recruiting chatbots. AI Multiple. Accessed on March 22 2022.
<https://research.aimultiple.com/recruiting-ai/>
- van Esch, P., Black, J.S. & Arli, D. (2021) Job candidates' reactions to AI-Enabled job application processes. *AI Ethics I*, 119–130. <https://doi.org/10.1007/s43681-020-00025-0>
- Fast, E., & Horvitz, E. (2017). Long-term trends in the public perception of Artificial Intelligence. *Proceedings of the AAAI Conference on Artificial Intelligence*, 31(1), 963–969. <https://doi.org/10.1609/aaai.v31i1.10635>
- Ferrer, X., van Neunen, T., Such, J. M., Cote, M., & Criado, N. (2020, August 11). *Bias and discrimination in AI: a cross-disciplinary perspective*. Cornell University.
<https://arxiv.org/pdf/2008.07309.pdf>
- Folger, N., Brosi, P., Stumpf-Wollersheim, Jutta, & Welp, I. (2022). Applicant reactions to digital selection methods: A signaling perspective on innovativeness and procedural justice. *Journal of Business and Psychology*. 37. 10.1007/s10869-021-

09770-3.

Gherheş V. (2018) Why are we afraid of artificial intelligence (Ai)? *European Review of Applied Sociology*, 11(17), 6-15.

Gilliland, S. (1993). The perceived fairness of selection systems: An organizational justice perspective. *The Academy of Management Review*, 18(4), 694-734.

Greszki, R., Meyer, M., & Schoen, H. (2015). Exploring the effects of removing "too fast" responses and respondents from web surveys. *Public Opinion Quarterly*, 79(2), 471-503.

Gonzalez, M. F., Capman, J. F., Oswald, F. L., Theys, E. R., & Tomczak, D. L. (2019). "Where's the I-O?" Artificial intelligence and machine learning in talent management systems. *Personnel Assessment and Decisions*, 5(3), 33-44.

Hair, J. F., Hult, G. T. M., Ringle, C. M., & Sarstedt, M. (2022). *A primer on partial least squares structural equation modeling (PLS-SEM)* (3 ed.). Sage.

Harris, M. M., Van Hove, G., & Lievens, F. (2003). Privacy and attitudes towards internet-based selection systems: A cross-cultural comparison. *International Journal of Selection and Assessment*, 11(2-3), 230-236.
<https://doi.org/10.1111/1468-2389.00246>.

Hess, R. A. (2022). *Applicant Reactions to Artificial Intelligence in Selection*. [Doctoral Dissertation, University of Georgia].

Highhouse, S. (2008). Stubborn reliance on intuition and subjectivity in employee selection. *Industrial and Organizational Psychology*, 1(33), 333-342.

HireVue. hirevue.com. (2023, May 25). <https://www.hirevue.com/>

Hsu, A. (2023, January 31). *Can bots discriminate?* it's a big question as companies use

AI for hiring. NPR. <https://www.npr.org/2023/01/31/1152652093/ai-artificial-intelligence-bot-hiring-eeoc-discrimination>

Huffcutt, A., van Iddekinge, C. H., & Roth, P. L. (2011). Understanding applicant behavior in employment interviews: A theoretical model of interviewee performance. *Human Resource Management Review*, 21(4), 353–367.
<https://doi.org/10.1016/j.hrmr.2011.05.003>

IBM.(2021, September). *The journey to AI and business-ready data begins with information ready architecture*. IBM.com.

<https://www.ibm.com/downloads/cas/LQDOQXEP>

Illinois General Assembly. (2019) Artificial intelligence video interview act, 820 ILCS 42. <https://www.ilga.gov/legislation/ilcs/ilcs3.asp?ActID=4015&ChapterID=68>

Interviewer.AI. (n.d.). *#1 End to end AI video interview platform*. Interviewer.AI.
<https://interviewer.ai>

Joshi, N. (2019, February 9). *Recruitment chatbots: is the hype worth it?* Forbes,
www.forbes.com/sites/cognitiveworld/2019/02/09/recruitment-chatbot-is-the-hype-worth-it/?sh=24316a34083f.

Kohn, L. S., & Dipboye, R. L. (1998). The effects of interview structure on recruiting outcomes. *Journal of Applied Social Psychology*, 28(9), 821–843.
<https://doi.org/10.1111/j.1559-1816.1998.tb01733.x>

Latham, G. P., & Finnegan, B. J. (1993). Perceived practicality of unstructured, patterned, and situational interviews. In H. Schuler, J. L. Farr, & M. Smith (Eds.), *Personnel selection and assessment: Individual and organizational perspectives* (pp. 41–55). Lawrence Erlbaum Associates, Inc.

- Leventhal, G.S., Karuza, J. and Fry, W.R. (1980) Beyond fairness: A theory of allocation preferences. *Justice and Social Interaction*, 3(1), 167-218.
- Lind, E. A., & Tyler, T. R. (1988). *The social psychology of procedural justice*. Plenum Press.
- Maurer, R. "Hirevue Discontinues Facial Analysis Screening." *SHRM*, 3 Feb. 2021, www.shrm.org/resourcesandtools/hr-topics/talent-acquisition/pages/hirevue-discontinues-facial-analysis-screening.aspx.
- McCarthy, J. Basic questions about. AI. *Stanford.edu*.
<http://jmc.stanford.edu/articles/whatisai/whatisai.pdf>
- McCarthy, J. M., Bauer, T. N., Truxillo, D. M., Anderson, N. R., Costa, A. C., & Ahmed, S. M. (2017). Applicant perspectives during selection: A review addressing “so what?,” “what’s new?,” and “where to next?” *Journal of Management*, 43(6), 1693–1725. <https://doi.org/10.1177/0149206316681846>
- Miller, D. T., & Ross, M. (1975). Self-serving biases in the attribution of causality: Fact or fiction? *Psychological Bulletin*, 82(2), 213–225.
<https://doi.org/10.1037/h0076486>
- Mitchell, T. M. (1997). Does Machine Learning Really Work?. *AI Magazine*, 18(3), 11.
<https://doi.org/10.1609/aimag.v18i3.1303>
- Petruzzello, G., Chiesa, R., Guglielmi, D., van der Heijden, B. I. J. M., de Jong, J. P., & Mariani, M. G. (2022). The development and validation of a multi-dimensional job interview self-efficacy scale. *Personality and Individual Differences*, 184, 111221. <https://doi.org/10.1016/j.paid.2021.111221>
- OpenAI. Retrieved May 2, 2024, from <https://chat.openai.com>

- Pew Research Center. (2021, April 7). *Mobile fact sheet*. Pew Research Center: Internet, Science & Tech. <https://www.pewresearch.org/internet/fact-sheet/mobile/>
- Ployhart, R. E., & Harold, C. M. (2004). The applicant attribution-reaction theory (AART): An integrative theory of applicant attributional processing. *International Journal of Selection and Assessment*, 12(1-2), 84–98.
<https://doi.org/10.1111/j.0965-075X.2004.00266.x>
- Ployhart, R. E., & Ryan, A. M. (1997). Toward an explanation of applicant reactions: An examination of organizational justice and attribution frameworks. *Organizational Behavior and Human Decision Processes*, 72(3), 308–335.
<https://doi.org/10.1006/obhd.1997.2742>
- Powell, D. M., Stanley, D. J., & Brown, K. N. (2018). Meta-analysis of the relation between interview anxiety and interview performance. *Canadian Journal of Behavioural Science / Revue canadienne des sciences du comportement*, 50(4), 195–207. <https://doi.org/10.1037/cbs0000108>
- Ryan, A. M., & Ployhart, R. E. (2000). Applicants' perceptions of selection procedures and decisions: A critical review and agenda for the future. *Journal of Management*, 26(3), 565–606. <https://doi.org/10.1177/014920630002600308>
- Schmitt, N., Oswald, F. L., Kim, B., Gillespie, M. A., & Ramsay, L. (2004). The impact of justice and self-serving bias explanations of the perceived fairness of different types of selection tests. *International Journal of Selection and Assessment*, 12(1–2), 160–171. <https://doi.org/10.1111/j.0965-075x.2004.00271.x>
- Schroth, H. A., & Pradhan Shah, P. (2000). Procedures: Do we really want to know them? An examination of the effects of procedural justice on self-esteem. *Journal*

of Applied Psychology, 85(3), 462–471. <https://doi.org/10.1037/0021-9010.85.3.462>

Shreve, J. T., Khanani, S. A., & Haddad, T. C. (2022). Artificial intelligence in oncology: Current capabilities, future opportunities, and ethical considerations. *American Society of Clinical Oncology Educational Book*, 42, 842–851. https://doi.org/10.1200/edbk_350652

Stahl, A. (2022, November 9). *How AI will impact the future of work and life*. Forbes. <https://www.forbes.com/sites/ashleystahl/2021/03/10/how-ai-will-impact-the-future-of-work-and-life/?sh=318b748c79a3>

Stone, D. L., Lukaszewski, K. M., Stone-Romero, E. F., & Johnson, T. L. (2013). Factors affecting the effectiveness and acceptance of electronic selection systems. *Human Resource Management Review*, 23(1), 50–70. <https://doi.org/10.1016/j.hrmr.2012.06.006>

Tabachnick, B. G., & Fidell, L. S. (2013). *Using multivariate statistics* (6th ed.). Pearson.

Tay, C., Ang, S., & Van Dyne, L. (2006). Personality, biographical characteristics, and job interview success: A longitudinal study of the mediating effects of interviewing self-efficacy and the moderating effects of internal locus of causality. *Journal of Applied Psychology*, 91(2), 446–454. <https://doi.org/10.1037/0021-9010.91.2.446>

Thibault, J. W., & Walker, L. (1975). *Procedural Justice: A Psychological Analysis*. L. Erlbaum Associates.

Thomas, M., & Powers, J. (2023, March 3). *The Future of AI: How Artificial Intelligence Will Change the World*. Built In. <https://builtin.com/artificial->

- Tippins, N., Oswald, F., & McPhail, S. (2021). Scientific, legal, and ethical concerns about AI-based personnel selection tools: A call to action. *Personnel Assessment and Decisions*, 7(2). 10.25035/pad.2021.02.001.
- Tross, S. A., & Maurer, T. J. (2008). The effect of coaching interviewees on subsequent interview performance in structured experience-based interviews. *Journal of Occupational and Organizational Psychology*, 81(4), 589–605.
<https://doi.org/10.1348/096317907X248653>
- Truxillo, D. M., & Bauer, T. N. (1999). Applicant reactions to test scores banding in entry-level and promotional contexts. *Journal of Applied Psychology*, 84(3), 322–339. <https://doi.org/10.1037/0021-9010.84.3.322>
- Walch, K. (2019, December 23). *Why cognitive technology may be a better term than artificial intelligence*. Forbes.
<https://www.forbes.com/sites/cognitiveworld/2019/12/22/why-cognitive-technology-may-be-a-better-term-than-artificial-intelligence/?sh=16c54030197c>
- Weiner, B. (1985). An attributional theory of achievement motivation and emotion. *Psychological Review*, 92(4), 548–573. <https://doi.org/10.1037/0033-295X.92.4.548>
- You, S., Nie, J., Suh, K., & Sundar, S. S. (2011). When the robot criticizes you...: Self-serving bias in human-robot interaction. *Proceedings of the Conference on Human-Robot Interaction*, 6(1), 295–296.
<https://doi.org/10.1145/1957656.1957778>
- Zhang, B. & Dafoe, A. (2019) “Artificial Intelligence Attitudes in America” *Center for*

Humanity Institute Oxford.

APPENDIX A

Tables

Table 1

Correlation Matrix

	Variable	Mean	SD	Process Fairness	Chance to Perform	Job Relatedness Predictive	Information Known	Consistency	Openness
1	Process Fairness	3.55	1.00	.909					
2	Chance to Perform	3.27	.96	.602**	.918				
3	Job Relatedness Predictive	3.23	1.01	.431**	.661**	.795			
4	Information Known	4.07	.76	.248**	.204*	-.006	.847		
	Variable	Mean	SD	Process Fairness	Chance to Perform	Job Relatedness Predictive	Information Known	Consistency	Openness
5	Consistency	3.78	.79	.192*	.223*	.059	.625**	.831	
6	Openness	3.79	.78	.450**	.469**	.264**	.463**	.605**	.866
7	Treatment	3.75	.81	.624**	.551**	.439**	.412**	.484**	.772**
8	Two-Way Communication	3.60	.85	.622**	.574**	.501**	.255**	.320**	.587**
9	Propriety of Questions	3.79	.74	.468**	.241**	.093	.471**	.576**	.557**
10	Job Relatedness Content	3.76	.88	.513**	.434**	.376**	.328**	.475**	.546**

Table 1 (continued)

	Variable	Mean	SD	Process Fairness	Chance to Perform	Job Relatedness Predictive	Information Known	Consistency	Openness
11	Self-Promotion SE	3.90	.72	.271**	.249**	.232*	.083	.208*	.158
12	Anxiety Management SE	3.75	.86	.277**	.311**	.231*	.070	.130	.162
13	Preparation SE	3.94	.71	.160	.151	.137	.194*	.288**	.217*
14	Logistical SE	4.26	.64	.181*	.108	.131	.248**	.241**	.143
15	Interaction SE	3.49	.90	.214*	.330**	.254**	.158	.185*	.161
16	Hours Employed	16.98	14.17	.099	.084	.106	-.167	-.063	.064
17	Hiring Mode Manipulation	.49	.50	-.208*	-.251**	-.193*	-.067	-.080	-.192*
18	Job Offer Manipulation	.50	.50	.177	.042	-.038	.095	.138	.150
19	White	.67	.47	-.103	-.067	-.176	.160	.029	-.006
20	Black	.17	.38	-.063	-.010	.078	-.259**	-.164	-.062

Table 1 (continued)

	Variables	Mean	SD	Process Fairness	Chance to Perform	Job Relatedness Predictive	Information Known	Consistency	Openness
21	Cohabiting/ Married/Engaged	.14	.35	.052	.171	.251**	.165	.221*	.192*
22	Non-Cohabiting	.24	.43	-.142	-.216*	-.279**	.043	-.107	-.116
23	Single	.63	.49	.089	.068	.065	-.156	-.064	-.035
24	Employed	.73	.44	.094	.090	.048	-.143	-.033	.036
25	Familiarity with AI	3.28	.98	.043	.040	-.041	.082	.037	.072
26	Has been interviewed before	.90	.30	.007	-.023	-.129	.020	-.067	.079
27	Expects to be interviewed in future	.98	.13	-.041	-.014	.111	.069	-.035	.027

Table 1 (continued)

	Variables	Treatmen t	Two-Way Communica tion	Propriety of Questions	Job Relatedness Content	Self- Promotion SE	Anxiety Managemen t SE	Preparation SE
7	Treatment	.894						
8	Two-Way Communication	.732**	.874					
9	Propriety of Questions	.556**	.431**	.714				
10	Job Relatedness Content	.559**	.635**	.509**	.843			
11	Self-Promotion SE	.257**	.341**	.193*	.231*	.854		
12	Anxiety Management SE	.236*	.301**	.180	.256**	.688**	.898	
13	Preparation SE	.197*	.270**	.243**	.193*	.524**	.497**	.843
14	Logistical SE	.308**	.304**	.194*	.162	.442**	.399**	.552**
15	Interaction SE	.253**	.362**	.164	.268**	.659**	.675**	.545**

Table 1 (continued)

	Variables	Treatment	Two-Way Communication	Propriety of Questions	Job Relatedness Content	Self- Promotion SE	Anxiety Management SE	Preparation SE
16	Hours Employed	.065	.138	.070	.094	.176	.140	-.141
17	Hiring Mode Manipulation	-.221*	-.234**	.004	-.081	-.141	-.067	-.041
18	Job Offer Manipulation	.108	.047	.161	.124	.006	-.067	.023
19	White	-.027	-.059	-.052	-.091	-.028	-.051	-.05
20	Black	-.120	.033	-.161	.001	-.004	.124	-.047
21	Cohabiting/ Married/Engaged	.172	.139	.158	.165	.112	.115	.125
22	Non-Cohabiting	-.183*	-.257**	-.065	-.077	-.106	-.077	-.087
23	Single	.038	.124	-.057	-.051	.013	-.014	-.013
24	Employed	.013	.096	.076	.117	.204*	.102	-.080
25	Familiarity with AI	.044	.024	-.014	-.021	.042	-.017	-.064
26	Has been interviewed before	.010	.021	.031	-.090	.079	.008	-.030
27	Expects to be interviewed in future	-.009	.000	.080	-.035	-.197*	-.151	.011

Table 1 (continued)

	Variables	Logistical SE	Interaction SE	Hours Employed	Hiring Mode Manipulat ion	Job Offer Manipulat ion	White Black	
14	Logistical SE	.731						
15	Interaction SE	.417**	.921					
16	Hours Employed	-.186*	.114					
17	Hiring Mode Manipulation	-.067	-.143	-.087				
18	Job Offer Manipulation	-.054	-.120	.050	.048			
19	White	.172	-.005	-.175	-.040	-.075		
20	Black	-.229*	-.058	-.012	.112	-.104	-.515**	
21	Cohabiting/ Married/Engaged	.044	.193*	.013	-.020	-.209*	.130	.005
22	Non-Cohabiting	-.064	-.153	.143	.139	.062	-.102	.051
23	Single	.025	-.004	-.136	-.107	.095	-.003	-.048
24	Employed	-.184*	.054	.741**	-.137	.091	-.146	-.018
25	Familiarity with AI	-.052	.033	.081	-.053	-.191*	-.072	.083
26	Has been interviewed before	-.018	.050	.039	.049	.000	.064	.004
27	Expects to be interviewed in future	.001	-.216*	-.152	-.130	.000	.048	-.113

Table 1 (continued)

	Variables	Cohabiting/ Married/ 'Engaged	Single	Employ ed	Familiarity with AI	Has been interviewed before	Expects to be interviewed in future
22	Non-Cohabiting	-.222*					
23	Single	-.518**	-.719**				
24	Employed	-.077	.120	-.051			
25	Familiarity with AI	.101	.074	-.136	.043		
26	Has been interviewed before	.132	.054	-.141	.173	.123	
27	Expects to be interviewed in future	.051	-.080	.033	-.077	-.160	-.042

Notes: Diagonal for scale variables represents Cronbach's Alpha values.

Hiring mode: 0 indicates traditional interview, 1 indicates AI interview.

Job offer: 0 indicates no job offer, 1 indicates job offer.

White: 0 indicates non-White, 1 indicates White.

Black: 0 indicates non-Black, 1 indicates Black.

Cohabiting/Married/Engaged: 1 indicates that the participant is either cohabiting married or engaged, 0 indicates they are not.

Non-Cohabiting: 0 indicates participant is not in a non-cohabiting relationship, 1 indicates that they are.

Single: 0 indicates the participant is in a relationship, 1 indicates they are single.

Employed: 0 indicates participant is not employed, 1 indicates that they are employed.

Has been interviewed before: 1 indicates that participants have been interviewed for a job before, 0 indicates they have not.

Expects to be interviewed in future: 1 indicates the participant expects to be interviewed sometime in the future, 0 indicates that they do not.

*. Correlation is significant at the 0.05 level (2-tailed).

**. Correlation is significant at the 0.01 level (2-tailed).

Table 2

Descriptive Statistics for Procedural Justice Dimensions

Variables	Pairwise n	Mean	Median	Standard Deviation	Minimum	Maximum	Skewness	Kurtosis
Process Fairness	122	3.55	3.67	1.010	1.00	5	-0.508	-0.449
Chance to Perform	122	3.27	3.25	0.967	1.00	5	-0.116	-0.559
Job Relatedness Predictive	123	3.23	3.00	1.000	1.00	5	-0.148	-0.251
Information Known	123	4.07	4.00	0.764	1.67	5	-0.490	-0.275
Consistency	124	3.78	4.00	0.785	2.00	5	-0.002	-0.886
Openness	123	3.79	3.75	0.776	1.00	5	-0.348	0.417
Treatment	117	3.75	3.80	0.813	1.00	5	-0.654	0.894
Two-Way Communication	121	3.60	3.80	0.852	1.60	5	-0.373	-0.296
Propriety of Questions	120	3.79	3.83	0.745	2.33	5	0.106	-0.909
Job Relatedness Content	123	3.76	4.00	0.876	1.50	5	-0.434	-0.160

Table 3

Descriptive Statistics for MJISE Dimensions

Variables	Pairwise n	Mean	Median	Standard Deviation	Minimum	Maximum	Skewness	Kurtosis
Self-Promotion SE	122	3.9016	4.00	0.72197	2.25	5	-0.107	-.721
Anxiety Management SE	123	3.7500	3.75	0.85711	1.25	5	-0.345	-.181
Preparation SE	124	3.9355	4.00	0.71416	2.25	5	-0.063	-.867
Logistical SE	123	4.2561	4.50	0.63654	2.75	5	-0.440	-.837
Interaction SE	123	3.4898	3.50	0.90105	1.00	5	-0.062	-.249

Table 4

Hypothesis Summary Table

Hypothesis	Alternative	Test	Result	Conclusion
1	AI would be rated less fair	MANOVA Regressions for Hypotheses 2, 3a, 3b, and 3c	Mixed support	AI was seen as less procedurally fair compared to traditional interviews on direct measure of PJ, but failed to find this using partial SPJS.
1 a-g	AI would be rated as less fair on SPJS dimensions	MANOVA	Unsupported	Unable to determine what PJ dimensions AI may violate if any
2	Job offers would be rated as fairer	Regression	Supported	Supports applicant self-serving bias
3a	Self-promotion SE moderates relationship between hiring mode and direct PJ	Hierarchical Regression	Unsupported	Self-promotion SE predicted fairness, but no interaction was found

3b	Anxiety management SE moderates relationship between hiring mode and direct PJ	Hierarchical Regression	Unsupported	Anxiety management SE predicted fairness, but no interaction was found
3c	Self-Promotion SE moderates relationship between hiring mode and direct PJ	Hierarchical Regression	Unsupported	Self-promotion SE predicted fairness, but no interaction was found

Hypothesis	Alternative	Test	Result	Conclusion
4a	Interaction SE moderates relationship between hiring mode and chance to perform	Hierarchical Regression	Unsupported	No interaction was found
4b	Anxiety management SE moderates relationship between hiring mode and chance to perform	Hierarchical Regression	Unsupported	Anxiety management SE predicted chance to perform but no interaction was found
4c	Interaction SE moderates relationship between hiring mode and chance to perform	Hierarchical Regression	Unsupported	No interaction was found

Table 5

Multivariate Tests for Predicting SPJS (Bauer et al, 2001) Dimensions

Effect	Pillai's Trace	F	p
Intercept	0.792	38.047*	<.001
Hiring Mode	0.134	1.552	.142
Job Offer	0.071	0.759	.654
White	0.130	1.493	.163
Black	0.089	0.977	.465
Cohabiting/ married/engaged	0.147	1.724	.095
Non-cohabiting	0.124	1.410	.196
Familiarity with AI	0.029	0.298	.974

Note * indicates significance at .01. level

Hiring Mode: 0 indicates traditional interview, 1 indicates AI interview

Job offer: 0 indicates no job offer, 1 indicates job offer.

Table 6

Regression Tables for Predicting SPJS (Bauer et al., 2001) Dimensions

Variable	R²	B	SE	t	p
Chance to perform	.177				.007
Intercept		3.435**	.380	9.048	<.001
AI Interview		-.485	.189	2.560	0.012
Job Offer		.266	.192	-1.385	.169
White		-.345	.231	-1.495	.138
Black		-.042	.290	-.145	.885
Cohabiting/Married/Engaged		.496	.272	1.828	.071
Non-Cohabiting		-.500	.224	-2.234	.028
Familiarity with AI		-.008	.096	-.085	.933
Job Relatedness Predictive	.228				<.001**
Intercept		3.925**	.392	10.013	<.001
AI Interview		-.361	.195	1.846	.068
Job Offer		.083	.199	-.419	.676
White		-.619	.238	-2.598	.011
Black		.004	.299	.013	.990
Cohabiting/Married/Engaged		.755*	.280	2.692	.008
Non-Cohabiting		-.607	.231	-2.629	.010
Familiarity with AI		-.125	.099	-1.258	.211

Table 6 (Continued)

Variable	R²	B	SE	t	p
Information Known	.132				.048
Intercept		4.015**	.301	13.330	<.001
AI Interview		-.237	.150	1.577	.118
Job Offer		.244	.153	-1.599	.113
White		.056	.183	.308	.759
Black		-.378	.230	-1.645	.103
Cohabiting/Married/Engaged		.458	.215	2.126	.036
Non-Cohabiting		.114	.178	.645	.521
Familiarity with AI		.009	.076	.121	.904
Consistency	.119				.079
Intercept		3.942**	.312	12.631	<.001
AI Interview		-.136	.156	.874	.384
Job Offer		.244	.158	-1.543	.126
White		-.171	.190	-.902	.369
Black		-.374	.238	-1.571	.119
Cohabiting/Married/Engaged		.560	.223	2.508	.014
Non-Cohabiting		-.124	.184	-.675	.501
Familiarity with AI		.014	.079	.179	.859
Openness	.168				.010
Intercept		3.782**	.304	12.446	<.001
AI Interview		-.336	.152	2.215	.029
Job Offer		.376	.154	-2.440	.016
White		-.135	.185	-.729	.468
Black		-.159	.232	-.686	.494
Cohabiting/Married/Engaged		.521	.217	2.397	.018

Table 6 (Continued)

Non-Cohabiting		-.214	.179	-1.198	.234
Familiarity with AI		.034	.077	.440	.661
Treatment	.181				.005
Intercept		3.983**	.319	12.496	<.001
AI Interview		-.414	.159	2.606	.011
Job Offer		.319	.162	-1.974	.051
White		-.251	.194	-1.295	.198
Black		-.279	.243	-1.148	.254
Cohabiting/Married/Engaged		.456	.228	1.998	.048
Non-Cohabiting		-.311	.188	-1.655	.101
Familiarity with AI		-.014	.081	-.174	.862
Two Way Communication	.161				.014
Intercept		3.800**	.329	11.542	<.001
AI Interview		-.461	.164	2.811	.006
Job Offer		.173	.167	-1.035	.303
White		-.247	.200	-1.235	.220
Black		.105	.251	.419	.676
Cohabiting/Married/Engaged		.289	.235	1.229	.222
Non-Cohabiting		-.415	.194	-2.141	.035
Familiarity with AI		-.034	.083	-.403	.688
Propriety of Questions	.113				.099
Intercept		4.195**	.291	14.429	<.001
AI Interview		-.039	.145	.267	.790
Job Offer		.188	.147	-1.278	.204
White		-0.384	0.177	-2.172	0.032
Black		-0.460	0.222	-2.070	0.041

Table 6 (Continued)

		0.426	0.208	2.047	0.043
		-0.102	0.171	-0.593	0.555
		-0.005	0.073	-0.066	0.948
Variable	R ²	B	SE	t	p
Job Relatedness Content	0.086				0.249
Intercept		4.068**	0.358	11.372	<.001
AI Interview		-0.165	0.178	0.925	0.357
Job Offer		0.210	0.181	-1.159	0.249
White		-0.391	0.217	-1.797	0.075
Black		-0.161	0.273	-0.591	0.556
Cohabiting/Married/Engaged		0.531	0.256	2.077	0.040
Non-Cohabiting		-0.068	0.211	-0.323	0.747
Familiarity with AI		-0.014	0.090	-0.158	0.875
Variable	R²	B	SE	t	p
Treatment	.181				.005
Intercept		3.983**	.319	12.496	<.001
AI Interview		-.414	.159	2.606	.011
Job Offer		.319	.162	-1.974	.051
White		-.251	.194	-1.295	.198
Black		-.279	.243	-1.148	.254
Cohabiting/Married/Engaged		.456	.228	1.998	.048
Non-Cohabiting		-.311	.188	-1.655	.101
Familiarity with AI		-.014	.081	-.174	.862
Two Way Communication	.161				.014
Intercept		3.800**	.329	11.542	<.001
AI Interview		-.461	.164	2.811	.006

Table 6 (Continued)

Variable	R²	B	SE	t	p
Job Offer		.173	.167	-1.035	.303
White		-.247	.200	-1.235	.220
Black		.105	.251	.419	.676
Cohabiting/Married/Engaged		.289	.235	1.229	.222
Non-Cohabiting		-.415	.194	-2.141	.035
Familiarity with AI		-.034	.083	-.403	.688
Propriety of Questions	.113				.099
Intercept		4.195**	.291	14.429	<.001
AI Interview		-.039	.145	.267	.790
Job Offer		.188	.147	-1.278	.204
White		-0.384	0.177	-2.172	0.032
Black		-0.460	0.222	-2.070	0.041
Cohabiting/Married/Engaged		0.426	0.208	2.047	0.043
Non-Cohabiting		-0.102	0.171	-0.593	0.555
Familiarity with AI		-0.005	0.073	-0.066	0.948
Job Relatedness Content	0.086				0.249
Intercept		4.068**	0.358	11.372	<.001
AI Interview		-0.165	0.178	0.925	0.357
Job Offer		0.210	0.181	-1.159	0.249
White		-0.391	0.217	-1.797	0.075
Black		-0.161	0.273	-0.591	0.556
Cohabiting/Married/Engaged		0.531	0.256	2.077	0.040
Non-Cohabiting		-0.068	0.211	-0.323	0.747
Familiarity with AI		-0.014	0.090	-0.158	0.875

Table 6 (Continued)

Variable	R ²	B	SE	t	p
Treatment	.181				.005
Intercept		3.983**	.319	12.496	<.001
AI Interview		-.414	.159	2.606	.011
Job Offer		.319	.162	-1.974	.051
White		-.251	.194	-1.295	.198
Black		-.279	.243	-1.148	.254
Cohabiting/Married/Engaged		.456	.228	1.998	.048
Non-Cohabiting		-.311	.188	-1.655	.101
Familiarity with AI		-.014	.081	-.174	.862
Two Way Communication	.161				.014
Intercept		3.800**	.329	11.542	<.001
AI Interview		-.461	.164	2.811	.006
Job Offer		.173	.167	-1.035	.303
White		-.247	.200	-1.235	.220
Black		.105	.251	.419	.676
Cohabiting/Married/Engaged		.289	.235	1.229	.222
Non-Cohabiting		-.415	.194	-2.141	.035
Familiarity with AI		-.034	.083	-.403	.688
Propriety of Questions	.113				.099
Intercept		4.195**	.291	14.429	<.001
AI Interview		-.039	.145	.267	.790
Job Offer		.188	.147	-1.278	.204
White		-0.384	0.177	-2.172	0.032
Black		-0.460	0.222	-2.070	0.041
Cohabiting/Married/Engaged		0.426	0.208	2.047	0.043

Table 6 (Continued)

Variable	R²	B	SE	t	p
Non-Cohabiting		-0.102	0.171	-0.593	0.555
Familiarity with AI		-0.005	0.073	-0.066	0.948
Job Relatedness Content	0.086				0.249
Intercept		4.068**	0.358	11.372	<.001
AI Interview		-0.165	0.178	0.925	0.357
Job Offer		0.210	0.181	-1.159	0.249
White		-0.391	0.217	-1.797	0.075
Black		-0.161	0.273	-0.591	0.556
Cohabiting/Married/Engaged		0.531	0.256	2.077	0.040
Non-Cohabiting		-0.068	0.211	-0.323	0.747
Familiarity with AI		-0.014	0.090	-0.158	0.875

Table 6 (Continued)

Note: Alphas are Bonferonni corrected: alpha/9

* indicates significance at .05/9 = .055 level

** indicates significance at .01/9 = .011 level

Table 7

Residual Normality

Analysis	Kolmogorov-Smirnov	Skewness	Kurtosis
Hypothesis 2	1.2*	-0.394	-0.625
Hypothesis 3a	0.092*	-0.604	-0.158
Hypothesis 3b	0.089*	-0.628	-0.103
Hypothesis 3c	0.085*	-0.564	-0.218
Hypothesis 4a	0.06	-0.253	-0.178
Hypothesis 4b	0.054	-0.326	-0.213
Hypothesis 4c	0.065	-0.236	-0.135

Note * indicates significance at .05. level

Table 8

Regression Table for Predicting Direct Measure of PJ (Truxillo & Bauer, 1999)

Variable	<i>B</i>	<i>t</i>	<i>p</i>
Constant	3.572**	23.967	<.001
Hiring Mode	-.437*	-2.467	.015
Job Offer	.377*	2.130	.035

Notes: $R^2 = .078$, $p = .008$

*. Correlation is significant at the 0.05 level

**. Correlation is significant at the 0.01 level

Hiring Mode: 0 indicates traditional interview, 1 indicates AI interview

Hypothesis 3 Tables

Table 9

Procedural Justice (Truxillo & Bauer, 1999) as a Function of Self-Promotion Self-Efficacy and Hiring Mode

Variable	R^2	ΔR^2	<i>B</i>	<i>t</i>	<i>p</i>
Model 1	0.104	0.104			
Constant			2.401**	4.737	<.001
Hiring Mode			-0.353*	-1.993	0.049
Self-Promotion SE			0.338**	2.742	0.007
Model 2	0.104	.000			
Constant			2.488**	3.560	<.001
Hiring Mode			-0.526	-0.539	0.591
Self-Promotion SE			0.316	1.836	0.069
Hiring mode $\times$ Self-Promotion SE			0.045	0.181	0.857

Notes: *. Correlation is significant at the 0.05 level

**. Correlation is significant at the 0.01 level

Hiring Mode: 0 indicates traditional interview, 1 indicates AI interview

Table 10

Procedural Justice (Truxillo & Bauer, 1999) as a Function of Anxiety-Management Self-Efficacy and Hiring Mode

Variable	R ²	ΔR ²	B	t	p
Model 1	0.109	0.109			
Constant			0.257**	6.316	<.001
Hiring Mode			-0.362*	-2.074	.040
Anxiety Management SE			0.307**	3.015	.003
Model 2	0.123	0.014			
Constant			2.953**	6.013	<.001
Hiring Mode			-1.467	-1.800	.074
Anxiety Management SE			0.206	1.647	.102
Hiring Mode × Anxiety Management			0.297	1.387	.168

Notes: *. Correlation is significant at the 0.05 level

**. Correlation is significant at the 0.01 level

Hiring Mode: 0 indicates traditional interview, 1 indicates AI interview

Table 11

Procedural Justice (Truxillo & Bauer, 1999) as a Function of Interaction Self-Efficacy and Hiring Mode

Variable	R ²	ΔR ²	B	t	p
Model 1	0.078	0.078			
Constant			3.003**	7.787	<.001
Hiring Mode			-0.368*	-2.038	0.044
Interaction SE			0.207*	2.049	0.043
Model 2	0.082	0.004			
Constant			2.754**	5.351	<.001
Hiring Mode			0.144	0.199	0.843
Interaction SE			0.276*	1.995	0.048
Hiring Mode × Interaction SE			-0.148	-0.730	0.467

Notes: *. Correlation is significant at the 0.05 level

**. Correlation is significant at the 0.01 level

Hiring Mode: 0 indicates traditional interview, 1 indicates AI interview

Hypothesis 4 Tables

Table 12

Chance to Perform (Bauer et al. 2001) as a Function of Self-Promotion Self-Efficacy and Hiring Mode

Variable	R ²	ΔR ²	B	t	p
Model 1	.194	.194			
Constant			2.983**	5.478	<.001
Hiring Mode			-.324	-1.906	.059
Self-Promotion			.272*	2.335	.021
Cohabiting/ Married/ Engaged			.127	.499	.619
Non-Cohabiting			-.307	-1.497	.137
Use AI Often			-.443	-.891	.375
Somewhat Familiar with AI			-.330	-.951	.344
Used AI			-.707*	-2.223	.028
Heard of AI			-.627	-1.969	.051
Model 2	.201	.007			
Constant			2.543**	3.620	<.001
Hiring Mode			.572	.623	.535
Self-Promotion			.390*	2.347	.021
Cohabiting/ Married/ Engaged			.117	.461	.645
Non-Cohabiting			-.304	-1.483	.141
Use AI Often			-.456	-.915	.362
Somewhat Familiar with AI			-.356	-1.024	.308
Used AI			-.740*	-2.315	.023
Heard of AI			-.655*	-2.048	.043
Hiring Mode × Self-Promotion			-.230	-.993	.323

Notes: *. Correlation is significant at the 0.05 level

**. Correlation is significant at the 0.01 level

Cohabiting/Married/Engaged: 1 indicates that the participant is either cohabiting married or engaged, 0 indicates they are not.

Non-Cohabiting: 0 indicates participant is not in a non-cohabiting relationship, 1 indicates that they are.

Hiring Mode: 0 indicates traditional interview, 1 indicates AI interview

Use AI Often (dummy code): 1 indicates respondent chose, "I am familiar with the (AI) technology and use it often"

Somewhat Familiar with AI (dummy code): 1 indicates respondent chose "I am somewhat familiar with AI and its uses but do not use it myself"

Used AI (dummy code): 1 indicates respondent chose, "I have used the (AI) technology before but am not an expert"

Heard of AI (dummy code): 1 indicates respondents chose, "I'd heard of AI but know very little about it"

A 0 on all AI familiarity (dummy code) items indicates the respondent chose, "I had never heard of AI technology before today"

Table 13

Chance to Perform (Bauer et al. 2001) as a Function of Anxiety-Management Self-Efficacy and Hiring Mode

Variable	R ²	ΔR ²	B	t	p
Model 1	.249	.249			
Constant			2.826**	6.226	<.001
Hiring Mode			-.381*	-2.332	.022
Anxiety Management			.338**	3.537	<.001
Cohabiting/ Married/ Engaged			.093	.374	.709
Non-Cohabiting			-.278	-1.402	.164
Use AI Often			-.474	-.978	.330
Somewhat Familiar with AI			-.310	-.908	.366
Used AI			-.789	-2.543	.012
Heard of AI			-.627	-2.021	.046
Model 2	.255	.006			
Constant			2.596**	5.053	<.001
Hiring Mode			.334	.436	.664
Anxiety Management			.408**	3.389	<.001
Cohabiting/ Married/ Engaged			.081	.328	.743
Non-Cohabiting			-.285	-1.439	.153
Use AI Often			-.485	-1.001	.319
Somewhat Familiar with AI			-.333	-.972	.333
Used AI			-.835**	-2.659	.009
Heard of AI			-.660*	-2.114	.037
Hiring Mode × Anxiety Management			-.191	-0.956	.341

Notes: *. Correlation is significant at the 0.05 level

**. Correlation is significant at the 0.01 level

Cohabiting/Married/Engaged: 1 indicates that the participant is either cohabiting married or engaged, 0 indicates they are not.

Non-Cohabiting: 0 indicates participant is not in a non-cohabiting relationship, 1 indicates that they are.

Hiring Mode: 0 indicates traditional interview, 1 indicates AI interview

Use AI Often (dummy code): 1 indicates respondent chose, "I am familiar with the (AI) technology and use it often"

Somewhat Familiar with AI (dummy code): 1 indicates respondent chose "I am somewhat familiar with AI and its uses but do not use it myself"

Used AI (dummy code): 1 indicates respondent chose, "I have used the (AI) technology before but am not an expert"

Heard of AI (dummy code): 1 indicates respondents chose, "I'd heard of AI but know very little about it"

A 0 on all AI familiarity (dummy code) items indicates the respondent chose, "I had never heard of AI technology before today"

Table 14

Chance to Perform (Bauer et al. 2001) as a Function of Interaction Self-Efficacy and Hiring Mode

Variable	R ²	ΔR ²	B	t	p
Model 1	.230	.230			
Constant			2.945**	6.534	<.001
Hiring Mode			-.325	-1.953	.053
Interaction SE			.302**	3.232	.002
Cohabiting/ Married/ Engaged			.070	.277	.782
Non-Cohabiting			-.245	-1.214	.227
Use AI Often			-.395	-.805	.422
Somewhat Familiar with AI			-.270	-.789	.432
Used AI			-.729*	-2.332	.022
Heard of AI			-.543	-1.733	.086
Model 2	.232	.002			
Constant			2.779**	5.056	<.001
Hiring Mode			-.012	.018	.985
Interaction SE			.349**	2.690	.008
Cohabiting/ Married/ Engaged			.075	.299	.766
Non-Cohabiting			-.241	-1.189	.237
Use AI Often			-.395	-.803	.424
Somewhat Familiar with AI			-.284	-.827	.410
Used AI			-.738*	-2.349	.021
Heard of AI			-.548	-1.743	.084
Hiring Mode × Interaction SE			-.097	-.530	.597

Notes: *. Correlation is significant at the 0.05 level

**. Correlation is significant at the 0.01 level

Cohabiting/Married/Engaged: 1 indicates that the participant is either cohabiting married or engaged, 0 indicates they are not.

Non-Cohabiting: 0 indicates participant is not in a non-cohabiting relationship, 1 indicates that they are.

Hiring Mode: 0 indicates traditional interview, 1 indicates AI interview

Use AI Often (dummy code): 1 indicates respondent chose, "I am familiar with the (AI) technology and use it often"

Somewhat Familiar with AI (dummy code): 1 indicates respondent chose "I am somewhat familiar with AI and its uses but do not use it myself"

Used AI (dummy code): 1 indicates respondent chose, "I have used the (AI) technology before but am not an expert"

Heard of AI (dummy code): 1 indicates respondents chose, "I'd heard of AI but know very little about it"

A 0 on all AI familiarity (dummy code) items indicates the respondent chose, "I had never heard of AI technology before today"

APPENDIX B

Figures

Figure 1 Hypothesized Moderation: Positive relationship for AI condition with non-cross interaction.

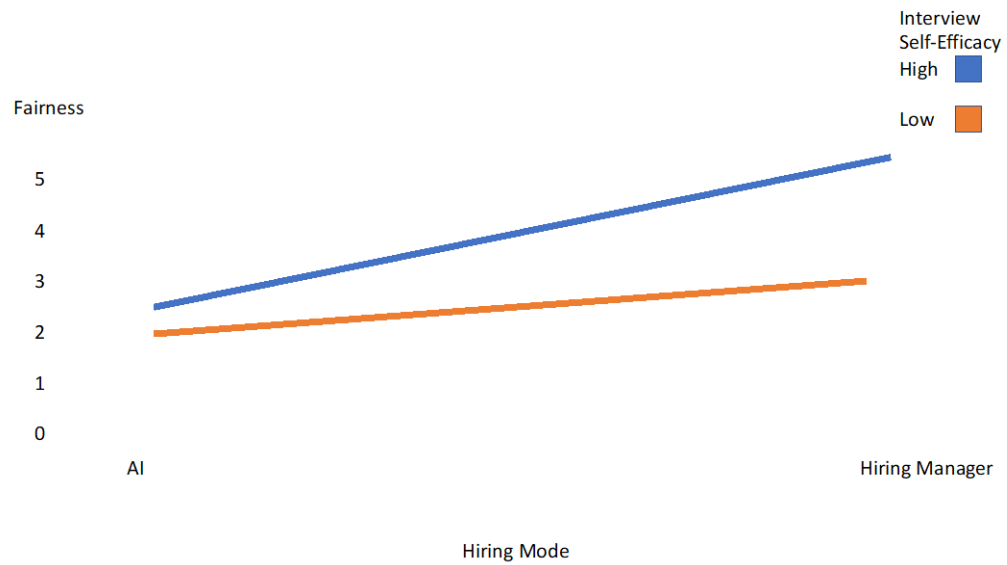
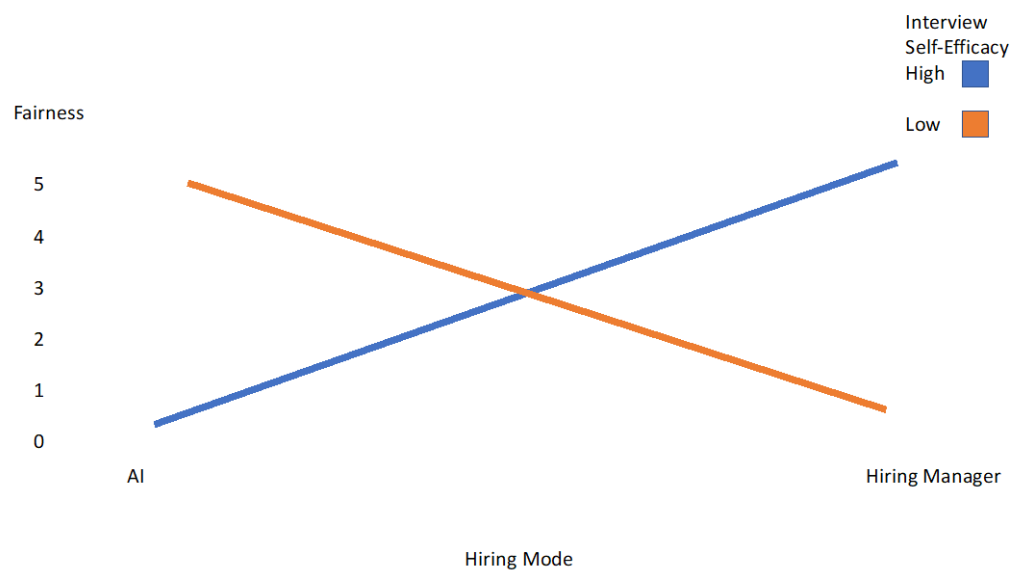


Figure 2 - Hypothesized Moderation: Negative relationship in AI condition with cross interaction



APPENDIX C

Vignette

PRE-INTERVIEW

You are applying for a job that you are interested in at Company XYZ. XYZ was founded twenty years ago. While the company is well known in your area, it is not particularly liked or disliked. You are qualified for the position, the salary and benefits meet your requirements, and you hear the working environment is acceptable. You have already submitted your resume and cover letter and you have been asked to complete an online video interview.

You then receive an email telling you what the interview process will be like. You are told how and when the interview will take place. It will be on a weekday in an afternoon in which you are available and take about 45 minutes. You will be video-interviewed by [a hiring manager/ an automated interview system powered by artificial intelligence that is similar to something like ChatGPT]. The [hiring manager/ artificial intelligence] has interviewed many people and [has training in conducting and appraising performance in job interviews/ is designed to conduct and appraise performance in job interviews]. For your interview, the [hiring manager/ artificial intelligence] will be speaking with you verbally over video chat. The interview will consist mostly of open ended questions. You will give your response verbally and the [hiring manager/ artificial intelligence] may choose to ask follow up questions.

You receive a pre-interview email which provides information about the testing format. You are given additional info about how the video call will work. You are told that you will need to connect to the video call with both audio and camera feed, and that

you can use a phone or computer as long as the device can support both of those requirements. There are instructions for those that do not have such a device, but please assume in this scenario that you do have a device that supports these requirements.

The email also contains a list of the upcoming interview questions. They include questions about your past work experience and your knowledge of the industry. The interview will also contain questions about your personality, skills, and ability to perform the tasks required. The email states that the [hiring manager/ artificial intelligence] will consider the content of your answers, as well as your tone of voice, facial expressions, and body language when appraising your performance in the interview. You will be asked what special skills and talents you have and what your biggest weaknesses are. The email also asks you to expect some problem solving tests during the interview.

Although you are not given the exact problem-solving questions, you are told that you will be given a problem to solve and 3 minutes to consider your answer. Then you will be asked to choose the best option from a list of five potential solutions. You will then be asked to defend your answer. You are told that the presented problems you will need to solve are common issues that occur in the role to which you are applying. After the interview, you will either be offered the job or rejected.

The pre-interview information email suggests that you refer any pre-interview questions to the company's support email (the address is provided). After the interview, you will have an opportunity to email questions or concerns to the company's support team.

INTERVIEW

You connect to the interview using your preferred device. The [hiring manager/

artificial intelligence] introduces [themselves/itself] and asks how your week has been. [They then ask/It then asks] if you have any hobbies. After hearing your answers the interviewer wishes you good luck on the interview before continuing. The interview proceeds first by asking the pre-provided test questions, with the problem solving portion at the end. All of the questions are exactly the ones described in the email. After answering this question, you move onto your problem solving questions. The problem solving section of the interview proceeds exactly as described in the email. The interviewer goes silent and waits 3 minutes before asking you what your answer to the problem is.

At the end of the interview you are told that before the final decision is made, you may address any questions to the online support email. You may also submit an appeal via email if you feel your performance on the test did not accurately reflect your true suitability for the role. For example, you may submit an appeal if circumstances outside of your control affected your performance. You may also submit an appeal if you believe an error has been made. You are told you will receive a decision with feedback in 2 to 5 business days.

Finally, the [hiring manager/ artificial intelligence] thanks you for your time and the call ends.

POST INTERVIEW

One week after your interview, you receive an email from the company containing information about the hiring decision.

Attached to the email is a report describing your performance in the interview, including scores for specific questions you answered. The email also provides the

company's final hiring decision.

[The email explains that you received a high score on the interview, and that the company will therefore extend an offer. You will be contacted with further information.]

[The email explains that you received a low score on the interview, and therefore will not receive an offer. Your application will be kept on file if other positions become available in the future. The email also contains contact information for a company representative, to be used if you believe an error has been made or would like to further discuss your interview results.]

APPENDIX D

Questionnaire

Informed Consent

I have read and the informed consent and agree to participate in this study.

Yes

No

I am 18 years or older

Yes

No

I am a US citizen

Yes

No

MJISE scale was presented

Can you see yourself interviewing for a job in the future?

Yes

No

In how many years do you think you will apply for a job? If you plan on applying

within a year, please answer 0.

Are you currently employed?

Yes

No

Are you employed part-time or full-time?

Part-time

Full-time

About how many hours do you spend in paid employment per a week? (Rounding up)

What is your age in years?

What is your gender?

Male

Female

Non-binary/third gender (please specify)

Prefer not to say

What is your ethnicity? (Choose more than one if applicable)

African American/Black

Native American

Latina/Hispanic

White

Middle-Eastern/Arab

Asian/Pacific Islander

Other (please specify)

Prefer not to answer

What is your relationship status?

Single

Separated

Married

Divorced

Widowed

Cohabiting with romantic partner

In a non-cohabiting relationship

Other (please specify)

Prefer not to say

Do you have children?

Yes

No

What is your current household income?

Under \$10,000

\$10,000 to \$20,000

\$20,001 to \$30,000

\$30,001 to \$40,000

\$40,001 to \$50,000

\$50,000 to \$60,000

\$60,001 to \$80,000

\$80,001 to \$100,000

Over \$100,001

What is the highest degree of education you have completed?

Less than high school

High school graduate

Some college

2 year degree

4 year degree

Professional/Master's degree

Doctorate

Vignette was presented

Truxillo and Bauer (1999) scale was presented

Bauer et al. scale

In the interview description, who was the interviewer?

An unnamed intern

An unnamed hiring manager

An artificial intelligence

The interview was cancelled/ there was no interview

There were multiple interviewers

The unnamed owner

I don't remember

Have you ever been interviewed for a job position before?

Yes

No

Have you ever been interviewed by an AI before? (If you are unsure select no)

Yes

No

How familiar are you with AI technology?

I'd never heard of AI technology before today

I'd heard of AI technology but know very little about it

I am somewhat familiar with AI technology and its uses but do not use it myself

I have used the technology before but am not an expert

I am familiar with the technology and use it often

I am an expert in AI technology